

JAO plan wykładu

1 III	Wstęp, wyrażenia regularne, automaty deterministyczne
8 III	Automaty niedeterministyczne, determinizacja, własności domknięcia
15 III	Równoważność automatów i wyrażeń regularnych, lemat o pompowaniu
22 III	Relacja Myhill-Neroda, minimalizacja
29 III	Problemy decyzyjne, monoidy
5 IV	Automaty dwukierunkowe, gramatyki, drzewa wyprowadzeń
12 IV	Lematy o pompowaniu, postaci normalne, automaty ze stosem
19 IV	brak wykładu
26 IV	Równoważność, deterministyczne automaty
3 V	przerwa
10 V	Algorytm CYK, Tw. Parikha
17 V	Maszyny Turinga, wielotaśmowe, determinizacja
24 V	KOLOKWIUM
31 V	Nierozstrzygalność (m.in. problemu stopu)
7 VI	Nierozstrzygalność ciąg dalszy, P vs NP, NP-zupełność SATa
14 VI	Klasy złożoności

Wykład 1

Technikalia

Moja strona, podstrona przedmiotu. Konsultacje piątek 14:15-15:45, też na maila. Bardzo ważne, żeby zadawać pytania. Zrobię ankietę po kilku wykładach.

Zasady zaliczania. Kolokwium: 40%

Egzamin: 60% - 20% test, 40% zadania

Egzamin poprawkowy: max(poprawkowy 100%, poprawkowy 60% + kolokwium 40%)

Do egzaminów dopuszczeni są wszyscy

Zadania z gwiazdką: 3 serie po około 2-4 zadania, każde zadanie to 100% do podziału po równo między tych co rozwiążą.

Mogę udostępnić nieformalne notatki, które robię dla siebie. Zrobić to?

Tematyka wykładu

Tematyka wykładu: po około $\frac{1}{3}$ automaty skończone, gramatyki i automaty ze stosem, maszyny Turinga. Po co ten przedmiot? Kilka odpowiedzi:

- zobaczyć modele obliczeń, do jest głównie przedmiot o modelach obliczeń przez różnorakie maszyny

- automaty skończone, co da się zrobić dysponując skończoną pamięcią, pojawia się to nieraz w informatyce
- maszyny Turinga - formalizują wszystko co może zrobić komputer
- na podstawie tego będzie złożoność obliczeniowa, która klasyfikuje które problemy są łatwe, a które nie dadzą się rozwiązać w rozsądnym czasie - na koniec zrobimy trochę tego
- techniki z teorii automatów przydają się w innych dziedzinach, ja miałem publikacje z dziedziny teorii baz danych i z dziedziny weryfikacji programów, które de facto bazowały na teorii automatów

Słowa

Dla nas słowo to będzie ciąg znaków z ustalonego (zwykle skończonego) alfabetu Σ . Wszystko z zasady można zapisać jako słowo, stąd pracujemy na słowach. Zwykle będzie nas interesować podział zbioru słów na dobre i niedobre. Czyli innymi słowy będziemy chcieli zdefiniować pewien zbiór słów. Zbiór słów nazywamy językiem. Przykłady z życia: poprawny numer PESEL, konta w banku, ale też zbiór ciągów 0-1 takich, że zakodowana liczba jest pierwsza albo zbiór wejść do algorytmu szukającego ścieżki z x do y takich, że algorytm odpowie TAK. Różne są języki, my na początek będziemy się zajmowali tymi prostszymi. Język odpowiada problemowi decyzyjnemu.

Słowo puste (0 liter) będziemy oznaczać przez epsilon. Długość słowa to liczba liter. Konkatenacja słów to napisanie jednego za drugim, oznaczamy kropką $\cdot$ lub $\cdot$, ale często piszemy po prostu uv . Definicja prefiksu, infiksu, sufiksu.

Wyrażenia regularne

Zdefiniujemy **wyrażenia regularne**, prosty sposób wyrażenia prostych języków. Będą one opisywały języki regularne.

Języki regularne mają bardzo fajne własności:

- prosta klasa zawierająca wiele naturalnych języków
- ma wiele alternatywnych definicji (typowa sytuacja w matematyce, która wskazuje na to, że mamy do czynienia z ważnym i właściwym pojęciem)
- zamknięte na wiele naturalnych operacji

Chciałoby się móc powiedzieć:

- nic nie należy do języka
- ustalone słowo należy do języka,
- możemy zapisać sumę dwóch zdefiniowanych języków, i bardziej skomplikowane
- jak zdefiniujemy języki K i L (tak zwykle oznaczamy języki) to umiemy też zdefiniować język KL , który składa się z dwóch połówek: pierwszej w K , drugiej w L
- jak zdefiniujemy język L , to umiemy też zdefiniować język L^* (gwiazdka Kleene'go), który składa się z wielu kawałków, z których każdy należy do L
- podobnie z językiem L^+

A więc formalnie języki regularne to (najmniejsza) klasa języków, która:

- zawiera zbiór pusty
- zawiera słowo epsilon
- zawiera słowo a dla każdej litery a w Σ
- jest zamknięta na sumę
- jest zamknięta na konkatenację
- jest zamknięta na gwiazdkę (Kleene'a)

Jeden ze sposobów zapisu (później poznamy fundamentalnie inne podejście) języków regularnych jest taki:

- zapisujemy sumę jako $+$
- konkatenację jako kropkę
- gwiazdkę jako $*$

Przykłady

a^* - zbiór słów złożonych z samych liter a

$(a+b)^*$ - zbiór słów złożonych z liter a i b

$abac$ - singleton słowa $abac$

$(ab)^*$ - zbiór słów, które są postaci ab ileś razy

$(ab)^*c$ - to co wyżej i jeszcze zakończone na c

$(ab)^* + b(ab)^* + (ba)^* + a(ba)^*$ - litery a i b na zmianę, albo $(a+eps)(ba)^*(b+eps)$

A jak zrobić: zaczynają się na a , kończą na b ?

$a(a+b)^*b$, a może $a(a+b+c)^*b$? To zależy od alfabetu, zawsze precyzujemy nad jakim alfabetem pracujemy. I piszemy wtedy $a\Sigma^*b$.

A jak powiedzieć: zbiór słów o ustalonym prefiksie / sufiksie / infiksie?

$w\Sigma^*$, Σ^*w , $\Sigma^*w\Sigma^*$

A pierwsza litera równa ostatniej? $a\Sigma^*a + b\Sigma^*b + c\Sigma^*c$

A jak zrobić: nie ma infiksu aa ? $(a + eps)(ba)^*(b + eps)$

A jak zrobić: nie ma infiksu aab ? To już będzie trudniej, chociaż da się.

Widać, że w ten sposób możemy zrobić wszystkie języki regularne po prostu konstruując wyrażenie.

Automaty skończone

Teraz przejdziemy do dużo bardziej naturalnego pojęcia, **automatów skończonych**.

Automaty skończone są innym sposobem opisu języków regularnych, dużo bardziej wygodnym.

Przykład: słowa o parzystej liczbie a i parzystej liczbie b .

Wyrażenie regularne: $(aa + bb + (ab + ba)(aa + bb)^*(ab + ba))^*$

Automat: łatwy do zrobienia. Pokazuję na przykładzie jak to wygląda.

Formalnie: deterministyczny automat skończony (deterministic finite automaton - DFA) $A = (Q, \Sigma, q_0, F, \delta)$ składa się ze skończonego zbioru stanów Q , jest nad skończonym alfabetem Σ , z wyróżnionym stanem początkowym q_0 , zbiorem stanów akceptujących F oraz funkcją przejścia $\delta: Q \times \Sigma \rightarrow Q$.

Bieg po $w = a_1 \dots a_n$ to ciąg tranzycji $q_i \xrightarrow{a_i} q_{i+1}$ takich, że q_1 jest początkowy. Bieg jest akceptujący jeśli $q_{n+1} \in F$. Automat akceptuje słowo w jeśli bieg po w jest akceptujący. Zbiór słów, które automat akceptuje to jego język, ozn. $L(A)$, mówimy, że A rozpoznaje $L(A)$.

Notacja: nieraz piszemy w takiej sytuacji $q_1 \xrightarrow{w} q_{n+1}$.

Idea jest taka, że automaty to po prostu coś ze skończoną pamięcią. Okazuje się, że języki, które są rozpoznawane przez DFA to dokładnie języki regularne.

Tw. Kleene'ego (1956)

Język jest rozpoznawalny przez DFA wtw. gdy jest regularny.

Dowód będzie na wykładzie 3.

Teraz zobaczmy kilka przykładów automatów.

Automaty dla:

- $(ab)^*$
- $(aaa)^*$
- pierwsza litera równa ostatniej (zrobiliśmy dotąd)
- parzystość $a =$ parzystość b
- jest podciąg aab
- nie ma podciągu aab

Wykład 2

Automaty niedeterministyczne

Teraz spojrzymy na nieco mocniejszy model, niedeterministyczny automat skończony (NFA).

Różnica jest taka, że on może mieć wiele stanów początkowych (może zacząć w jednym z wielu miejsc) i może mieć wiele tranzycji z tego samego stanu po tej samej literze.

Przykład: Σ^*abc

Formalnie: $A = (Q, \Sigma, I, F, \delta)$, gdzie

Q - skończony zbiór stanów

Σ - skończony alfabet

$I \subseteq Q$ - stany początkowe

$F \subseteq Q$ - stany akceptujące

δ - **relacja** przejścia, $\delta \subseteq Q \times \Sigma \times Q$

Bieg - jak poprzednio ciąg tranzycji, zaczyna się w stanie początkowym

Bieg jest akceptujący jeśli kończy się w stanie końcowym.

Teraz: po jednym słowie może być wiele biegów.

Słowo jest akceptowane przez automat jeśli istnieje (choćby jeden) bieg akceptujący.

Język A to zbiór słów akceptowanych.

Powiemy, że automaty są równoważne jeśli mają te same języki.

Pytanie: czy dla każdego automatu niedeterministycznego istnieje równoważny automat deterministyczny? Ankieta studentów. Odpowiedź: TAK.

Determinizacja

Tw.

Dla każdego automatu niedeterministycznego istnieje równoważny automat deterministyczny.

Dowód

Weźmy NFA $A = (Q, \Sigma, I, F, \delta)$

Zrobimy równoważny DFA $A' = (Q', \Sigma, q_0', F', \delta')$

Idea ogólna DFA jest taka, że po prostu pamiętamy wszystko to, co jest nam potrzebne o słowie, które wczytaliśmy do tej pory i dla każdego stanu pamięci mamy jeden stan.

Idea: patrzemy gdzie może A pójść po słowie w, do których stanów i to pamiętamy. Czyli innymi słowy pamiętamy podzbiór stanów.

$$Q' = P(Q)$$

$$q_0' = \{I\}$$

$$F' = \{S \mid S \cap F \neq \emptyset\}$$

$$\delta'(S, a) = \{q \mid \exists p \in S \text{ } (p, a, q) \in \delta\}$$

Dlaczego to jest ok. Pokażmy, że istotnie $L(A) = L(A')$. Inkluzja $L(A) \subseteq L(A')$ - jeśli jest bieg po w, to po tym w idziemy z I do jakiegoś S i w tym S jest stan akceptujący. Inkluzja $L(A') \subseteq L(A)$ - jeśli jest bieg po w, to znaczy, że z pewnego stanu z I do się dojść po w do pewnego stanu z F, koniec dowodu. Komentarz: w zasadzie to jest masło maślane, widać, że A' dobrze symuluje A.

c.n.d.

Przykład dla $\Sigma^* abc$

Widać, że wszystkie stany osiągalne A' będą zawierały stan początkowy A, więc wystarczy rozważyć ich 8.

Uwaga: determinizacja może być wykładnicza. Przykład: n-ta litera od końca to a.

NFA może to zrobić używając $n+1$ stanów. **Pytanie:** ile potrzebuje DFA?

Co najmniej 2^n . Formalnie będziemy umieli to zrobić za 1-2 wykłady, ale teraz zobaczmy z grubsza dlaczego. Idea: musi pamiętać ostatnie n liter, bo gdyby coś zapomniał, to mogłoby się okazać, że to właśnie jest ważne. Czyli potrzeba 2^n stanów, no i tyle wystarczy.

Operacje na językach

Będziemy wymiennie używali automatów deterministycznych i niedeterministycznych, bo wiemy, że ich wyrażalność jest taka sama. Natomiast nie używamy na razie Tw. Kleene'go o tym, że automaty (DFA lub NFA) mają tę samą wyrażalność co wyrażenia regularne, bo tego żeśmy jeszcze nie udowodnili.

Dopełnienie - jak?

Dla automatu deterministycznego dużo łatwiej. Po prostu zamieniamy stany akceptujące i nieakceptujące. Czy to działa? Tak po prostu to nie, musimy mieć własność taką, że po każdym słowie jest bieg. Więc wprowadzamy dodatkowy stan - "śmietnik", do którego idą wszystkie tranzycje, które wcześniej nie istniały. Po tej modyfikacji funkcja przejścia automatu jest już prawdziwą funkcją, a nie tylko funkcją częściową. W konstrukcji dopełnienia ten stan "śmietnik" stanie się akceptujący. W przyszłości będziemy używać automatów ze śmietnikiem lub bez, w zależności od wygody. Zwykle wygodnie jest go nie rysować, ale czasem, jak przy dopełnieniu, warto mieć własność taką, że delta jest funkcją, wtedy go używamy.

Tu ciekawym pytaniem jest: jak duży może być NFA dla $\bar{L}$ (dopełnienia L) w stosunku do NFA dla L. Bo dla DFA są takie same. Pytanie: jak automat ma n stanów to jak długie może być najkrótsze słowo, które akceptuje (o ile akceptuje jakiś)? Oczywiście nie dłuższe niż n-1. Okazuje się, że da się zrobić taki NFA A, że najkrótsze słowo, którego nie akceptuje jest wykładniczej długości względem n. Wtedy w dopełnieniu $L(A)$ najkrótsze akceptowane słowo jest wykładniczej długości względem n, czyli automat musi mieć przynajmniej taką wielkość.

Suma, przecięcie

Suma - robimy automat $A = (Q, \Sigma, q, F, \delta)$ z automatów $A_1 = (Q_1, \Sigma, I_1, F_1, \delta_1)$ oraz $A_2 = (Q_2, \Sigma, I_2, F_2, \delta_2)$. Po prostu A będzie rozłączną sumą A_1 i A_2 , czyli formalnie $Q = Q_1 \cup Q_2$, $I = I_1 \cup I_2$, $F = F_1 \cup F_2$, $\delta = \delta_1 \cup \delta_2$.

Przecięcie - nie da się już tak łatwo. Idea jest taka, że będziemy symulowali oba automaty i akceptowali jeśli oba zaakceptowały.

Jak wyżej robimy A z A_1 i A_2 . Ustalamy:

- $Q = Q_1 \times Q_2$ (nowy stan musi pamiętać oba stany)
- $I = I_1 \times I_2$
- $F = F_1 \times F_2$
- δ też zachowuje się produktowo, tzn. $\delta = \{(p_1, p_2) \xrightarrow{a} (q_1, q_2) \mid p_1 \xrightarrow{a} q_1 \text{ oraz } p_2 \xrightarrow{a} q_2\}$

Łatwo zauważyć, że $w \in L(A)$ wtw. $w \in L(A_1)$ oraz $w \in L(A_2)$, można pokazać implikację w dwie strony.

Zauważmy, że podobnie można też zrobić konstrukcję dla sumy, tyle, że wtedy dajemy $F = (F_1 \times Q_2) \cup (Q_1 \times F_2)$, czyli akceptujemy jeśli co najmniej jeden ze stanów jest akceptujący.

Automaty z epsilon-przejdściami

Wprowadzimy teraz nowy model automatu (ale nie bardzo różny), który będzie przydatny do różnych technicznych rozważań. Okazuje się, że będzie miał tę samą wyrażalność co DFA i NFA. Najpierw przykład: $(ab)^* (aaa)^*$ (dwa automaty połączone epsilon-przejdciem).

Automat z epsilon-przejdściami to po prostu NFA, który może mieć puste słowo na tranzycjach.

Jego język jest zdefiniowany normalnie, tyle, że słowo na ścieżce po prostu czyta epsilony.

Oczywiście automaty z epsilon-przejdściami akceptują wszystkie języki akceptowane przez NFA, bo po prostu każdy NFA jest też automatem z epsilon-przejdściami (tyle, że zeroma).

Pytanie jak w drugą stronę, czy mogą coś więcej? **Pytanie** do sali. Odpowiedź: NIE.

Tw.

Automaty z epsilon-przejdściami akceptują te same języki co automaty niedeterministyczne (i deterministyczne też zresztą).

Dowód

Przerabiamy automat z eps-przejdściami na automat równoważny (o tym samym języku), ale taki, który nie ma już eps-przejdść. Tzn. eliminujemy eps-przejdścia. Najpierw dla każdej ścieżki typu $p_1 \xrightarrow{\text{eps}} p_2 \xrightarrow{\text{eps}} \dots \xrightarrow{\text{eps}} p_k \xrightarrow{a} q_1 \xrightarrow{\text{eps}} \dots \xrightarrow{\text{eps}} q_m$ dodajemy nową tranzycję $p_1 \xrightarrow{a} q_m$. Nowe przejdścia nie zmieniają języka. Dla każdego akceptowanego niepustego słowa znajdziemy bieg, który nie używa eps-przejdść. Więc prawie możemy teraz usunąć eps-przejdścia, trzeba tylko uważać na puste słowo. No to usuwamy eps-przejdścia i patrzymy czy zmieniła się sytuacja pustego słowa. Jeśli tak, to dodajemy je sztucznie (dodając nowy stan równocześnie początkowy i akceptujący).

c.n.d.

Wykład 3

Dowód tw. Kleene'go

Jesteśmy już gotowi, żeby udowodnić równoważność wyrażeń regularnych i automatów.

Przypomnijmy, że to właśnie mówi twierdzenie Kleene'go (1956).

Potrzeba pokazać, że dla każdego wyrażenia regularnego r istnieje automat A taki, że $L(A) = L(r)$ oraz odwrotnie, że dla każdego automatu A istnieje wyrażenie regularne r takie, że $L(r) = L(A)$.

Zacniemy od przerabiania wyrażenia na automat, czyli pokazywania, że dla każdego wyrażenia r istnieje automat A taki, że $L(A) = L(r)$. Będziemy to pokazywać indukcyjnie po budowie wyrażenia regularnego. Czyli zacniemy od wyrażeń bazowych, tj.: puste, epsilon, a , a potem będziemy pokazywać dla coraz większych wyrażeń. Pokażemy innymi słowy automaty dla wyrażeń: puste, epsilon, a , a potem pokażemy, że jeśli mamy dla wyrażeń r_1 i r_2 mamy odpowiednio automaty A_1 i A_2 , to umiemy zbudować automat dla wyrażeń r_1+r_2 , $r_1 r_2$, r_1^* . W ten sposób możemy składać z małych wyrażeń coraz większe, czyli formalnie właśnie dowodzić indukcyjnie po budowie wyrażenia.

Automaty dla wyrażeń: puste, epsilon, a są banalne (narysować).

Zajmijmy się sumą. Będziemy robić automat z epsilon-tranzycjami dla $L(A1) \cup L(A2)$. **Pytanie** do wszystkich - jak to zrobić? Łatwo - robię nowy stan początkowy i z niego eps-tranzycje do wszystkich stanów początkowym w A1 i w A2 (czyli odpowiednio z I1 i z I2).

A jak dla konkatenacji, czyli dla $L(A1) L(A2)$? Też łatwo - dodaję eps-tranzycje pomiędzy każdym stanem akceptującym w A1 (czyli z F1) a każdym stanem początkowym w A2 (czyli z I2). Pokazać, że to istotnie to samo.

A dla $L(A1)^*$? Dodaję eps-tranzycje z każdego stanu akceptującego (z F1) do każdego stanu początkowego (z I1). Czy to jest ok - **pytanie** do sali. Nie, ale prawie. W $L(A1)^*$ powinno być słowo puste, ta konstrukcja niekoniecznie o to dba. Dodajemy więc niezależnie słowo puste (nowy stan akceptujący i eps-przejdzie do niego).

To kończy dowód w jedną (nieco prostszą) stronę.

Zasymulujmy jak to działa na przykładzie wyrażenia $(a+b+c)^*c(a+b)^* + (a+c)^*b(a+b+c)^*$.

Teraz robimy w drugą stronę, mianowicie dla danego automatu budujemy równoważne wyrażenie. Na początek zakładamy, że nasz automat ma dokładnie jeden stan akceptujący, język każdego automatu deterministycznego jest skończoną sumą takich właśnie automatów, po prostu w każdym z automatów zostawiamy jako akceptujący tylko jeden ze stanów akceptujących w oryginalnym automacie. Nasz plan postępowania jest następujący: bierzemy automat i będziemy po kolei zmniejszać liczbę jego stanów pisząc na krawędziach nie tylko litery, ale również wyrażenia. Na koniec dostaniemy automat z dwoma stanami i jakimś wyrażeniem na krawędzi, to będzie interesujące nas wyrażenie albo z jednym stanem (i wtedy szukane wyrażenia to gwiazdka z wyrażenia na pętli). Formalnie rzecz biorąc wprowadzamy tutaj znów nowy model automatu, tj. automat, który na krawędziach ma napisane nie tylko litery albo słowa puste, ale wyrażenia. Wtedy przejście jest po dowolnym słowie należącym do języka wyrażenia. Nie będziemy robić tego formalnie, bo właśnie powiedzieliśmy precyzyjnie o co chodzi, a nie chodzi o formalność, tylko o precyzję. **Pytanie**: czy ktoś chce, żeby to napisać nieco dokładniej?

No to teraz jak przerabiamy? Zrobić to na przykładzie automatu dla: parzysta liczba a i parzysta liczba b. Wychodzi nieco brzydko, ale da się zasymulować. W ogólności: jeśli eliminujemy stan q, to dla dowolnych stanów p1, p2 jeśli wcześniej było:

- $p1 \xrightarrow{r1} q$
- $q \xrightarrow{r2} p2$
- $q \xrightarrow{r3} q$
- $p1 \xrightarrow{r4} p2$

To teraz piszemy $p1 \rightarrow (r4 + r1(r3)*r2) \rightarrow p2$. No bo tak właśnie mogliśmy dojść z $p1$ do $p2$ (nie przechodząc przez q , bądź przechodząc przez nie). Pokazujemy, że jeden krok tej modyfikacji nie zmienia języka, co kończy dowód.

c.n.d.

Lemat o pompowaniu

Teraz już wiemy, że wszystkie automaty i wyrażenia, które badaliśmy opisują tę samą klasę języków: języki regularne. Fajnie byłoby coś zrozumieć o tych językach. W szczególności dla ustalonego języka dobrze byłoby umieć pokazać, że on nie jest regularny. Mówiliśmy już, że intuicyjnie $L = \{w \mid \#a(w) = \#b(w)\}$ nie jest regularny (bo trzeba by pamiętać różnicę pomiędzy liczbą a , a liczbą b do tej pory, a to może być dowolnie duża), ale chcemy to umieć sformalizować. Teraz stworzymy pierwsze narzędzie do tego: lemat o pompowaniu. Niedługo nauczymy się drugiego (moim zdaniem wygodniejszego) narzędzia, ale zacznijmy od tego.

Lemat o pompowaniu ma to do siebie, że łatwiej zapamiętać jego dowód niż jego sformułowanie. Zresztą (moim zdaniem) dowód jest też bardziej interesujący niż sformułowanie. Dlatego polecam zapamiętania dowodu, bo wtedy łatwiej odtworzyć treść lematu.

Założmy, że mamy automat deterministyczny o n stanach.

Najpierw obserwacja: jeśli mamy bieg po słowie w w długości co najmniej n , to odwiedza on przynajmniej $n+1$ stanów, więc z zasady szufladkowej Dirichleta któryś stan zostanie odwiedzony przynajmniej 2 razy. A więc możemy przedstawić nasz bieg w postaci:

$$p \xrightarrow{s} q \xrightarrow{t} q \xrightarrow{u} r,$$

gdzie $w = stu$. No ale pętlę $q \xrightarrow{t} q$ możemy powtarzać wiele razy, więc dowolne słowo postaci $st^k u$ również będzie miało bieg akceptujący, czyli będzie należało do języka. Co więcej możemy wybrać słowa s i t tak, żeby $|st| \leq n$, bo już na prefiksie słowa w o długości n powtórzy się jakiś stan q .

To już właściwie koniec dowodu, ale najpierw sformułujmy lemat.

Lemat o pompowaniu

Dla każdego języka regularnego L istnieje n w $\mathbb{N}$ takie, że dla każdego słowa w o długości co najmniej n istnieją słowa s, t, u takie, że:

- $w = stu$
- $|st| \leq n, t \neq \epsilon$
- $st^k u \in L$ dla dowolnego $k \in \mathbb{N}$

Dowód

Bierzemy n równe liczbie stanów automatu deterministycznego (choć to niekonieczne) dla L . Wtedy lemat wynika z obserwacji powyżej.

c.n.d.

Przykłady zastosowania lematu

Przykład 1: rozważmy język $L = \{w \mid \#a(w) = \#b(w)\}$. Przypuśćmy, że jest regularny.

Wówczas istnieje dla niego pewna liczba n z lematu o pompowaniu.

Weźmy $w = a^n b^n \in L$. Wtedy $w = st^i u$. Ponieważ $|st^i| \leq n$, to musi być $t = a^i$, $i > 0$.

Zatem $st^{2i}u = a^{n+i} b^n$ też należy do L (z lematu). No ale ewidentnie $a^{n+i} b^n$ nie należy do L , sprzeczność.

Przykład 2: niech $L = \{a^i b^j \mid i \geq j\}$. Przypuśćmy, że L jest regularny. Wówczas istnieje dla niego n z lematu o pompowaniu. Weźmy $w = a^n b^n \in L$. Wtedy $w = st^i u$, podobnie jak wcześniej mamy $t = a^i$, $i > 0$. Ale jak zbadamy $st^{2i}u$, to on naprawdę należy do L . Możemy jednak wziąć $k = 0$. I wtedy mamy, że $su = a^{n-i} b^n$ też należy do L , ale jednak nie należy. Sprzeczność.

Przykład 3: niech L = język palindromów. Bierzemy $w = a^n b a^n$, pompujemy, działa.

Przykład 4: niech $L = \{w \mid \#a(w) = \#b(w) = 2\#c(w)\}$. Można ręcznie. Ale można też: gdyby był automat dla L , to automat, w którym wszystkie c zamieniamy na epsilon rozpoznawałby $\{w \mid \#a(w) = \#b(w)\}$, a ten już wiemy, że jest nieregularny.

Przykład 5: niech $L = \{a^i b^j c^k d^m \mid i + j = k + l\}$. Języki regularne są zamknięte na przecięcie, więc przetnijmy L z a^*c^* . Jeśli tamten był regularny, to nowy też będzie. Ale $L \cap L(a^*c^*) = \{a^i c^i\}$, który jest nieregularny jak łatwo pokazujemy używając lematu o pompowaniu.

Uwaga: to nie jest tak, że język jest regularny wtedy i tylko wtedy, gdy spełnia lemat o pompowaniu. Są języki nieregularne, które spełniają lemat o pompowaniu.

Przykładem może być język: $L = \{c a^n b^n\} \cup \{c^k w \mid w \text{ zaczyna się nie na } c, k \neq 1\}$.

Druga część jest regularna, więc może być napompowana, a pierwsza część może być napompowana przez napompowanie lub odpompowanie c i wpadnięcie do drugiej części. Czyli L spełnia lemat o pompowaniu chociaż nie jest regularny (co się okaże później, ale intuicyjnie już widać).

Uwaga: można sobie dowodzić własne wersje lematu. Dużo bardziej należy z tego zapamiętać technikę usuwania/dodawania kawałków pomiędzy równymi konfiguracjami niż sformułowania lematu!

Wykład 4

Relacja Myhill-Nerode (John Myhill, Anil Nerode)

Teraz zajmiemy się pojęciem, które da nam inne narzędzie do sprawdzania, czy język jest regularny. Poza tym okaże się przydatne np. do zmniejszania rozmiaru automatu

deterministycznego przy zachowaniu jego języka. Pomysł bierze się z następującego faktu.

Ustalmy automat deterministyczny A , niech $L = L(A)$. Powiedzmy, że dwa słowa: u i v prowadzą do tego samego stanu q w automacie A . Wówczas mają one tę samą przyszłość, to znaczy uw

$w \in L \Leftrightarrow vw \in L$, bo to po prostu znaczy, że słowo w przechodzi ze stanu q do jakiegoś stanu akceptującego. To motywuje następującą definicję. Dla ustalonego języka L (niekoniecznie regularnego) relacja równoważności $\sim_L$ na zbiorze słów Σ^* jest zdefiniowana jako:

$u \sim_L v$ wtw. dla każdego słowa w zachodzi $uw \in L \Leftrightarrow vw \in L$

Mówimy o niej relacja Myhill-Neroda. Intuicyjnie znaczy to, że u i v mają te same przyszłości. Uwaga: inaczej można napisać, że $u^{-1}L = v^{-1}L$, to oznacza to samo. Formalnie $u^{-1}L$ nazywa się ilorazem lewostronnym języka. Łatwo zauważyć, że $\sim_L$ jest relacją równoważności. Jej istota ujęta jest w następującym twierdzeniu.

Zakładamy, że funkcja przejścia automatu musi być funkcją całkowitą. Ta definicja jest jednak lepsza. Nieraz rysujemy automaty bez stanu śmietnik, ale formalnie powinien on tam być.

Tw. Myhill-Neroda (1958)

L jest regularny wtw. relacja $\sim_L$ ma skończenie wiele klas abstrakcji (tzn. skończony indeks)

Dowód:

Najpierw pokażemy implikację w prawo. Jest łatwa. Niech A to DFA taki, że $L = L(A)$.

Jak powiedzieliśmy wcześniej jeśli biegi po u i v prowadzą do tego samego stanu w A , to $u \sim_L v$. Czyli klas abstrakcji jest co najwyżej tyle, co stanów A , czyli skończenie wiele.

Teraz pokażemy implikację w lewo. Załóżmy, że $\sim_L$ ma skończenie wiele klas abstrakcji. Zauważmy, że jeśli $u \sim_L v$, to także $ua \sim_L va$ dla dowolnej litery a . Zdefiniujemy następujący automat.

Stanami są klasy abstrakcji $\sim_L$. Przejścia są zdefiniowane jako $[u] \xrightarrow{a} [ua]$. Jak pokazaliśmy powyżej to nie zależy od wyboru reprezentanta klasy (czyli u), czyli automat jest deterministyczny. Stan początkowy to $[\epsilon]$, a stany końcowe to $[u]$ dla $u \in L$. Zauważmy, że w klasach akceptujących wszystkie słowa są akceptujące.

Teraz zobaczymy, że ten automat rozpoznaje język L . Dla dowolnego słowa u dojdziemy w nim do $[u]$. Jeśli $u \in L$, to ten stan jest akceptujący, a wpp. nie, czyli rzeczywiście dokładnie słowa z L zostaną zaakceptowane.

c.n.d.

Pytanie: czy to samo by zachodziło dla ilorazów prawostronnych języka L ? Tak, tu wszystko jest symetryczne. L regularny $\Leftrightarrow \text{reverse}(L)$ regularny.

Przykłady

Łatwo teraz pokazywać używając tego narzędzia, że dany język nie jest regularny.

Przykład 1: $\{a^n b^n\}$. Słowa a^i są nie w relacji M-N. Zobaczmy, że a^i oraz a^j nie są w relacji. Dla sufiksu a^i mamy $a^i b^i \in L$, ale $a^j b^i \notin L$.

Przykład 2: $\{w \mid \#a(w) = \#b(w)\}$. Ten sam przykład.

Przykład 3: weźmy przykład języka nieregularnego, dla którego działał lemat o pompowaniu, czyli $L = \{ca^n b^n\} \cup \{c^k w \mid w \text{ zaczyna się nie na } c, k \neq 1\}$. Pokażemy, że słowa ca^i dla $i > 0$ są nie w relacji. Faktycznie, ca^i oraz ca^j , dla $i \neq j$, rozróżniają się sufiksem b^i .

Automat minimalny

Relacja M-N jest też mocno związana z minimalnym automatem deterministycznym. Zauważmy, że pokazaliśmy, że jeśli A to DFA oraz $L = L(A)$, to liczba stanów $\geq \text{indeks}(\sim_L)$. Z drugiej strony pokazaliśmy, że da się zbudować DFA A' taki, że $L(A') = L(A)$ oraz liczba stanów $A' = \text{indeks}(\sim_L)$. Czyli innymi słowy pokazaliśmy, że najmniejszy pod względem liczby stanów automat deterministyczny dla L ma ich dokładnie $\text{indeks}(\sim_L)$. To odpowiada intuicji: to jest dokładnie ta informacja, którą musimy pamiętać.

Fakt

Minimalny automat deterministyczny dla L ma $\text{indeks}(\sim_L)$ stanów.

Zauważmy, że łatwo w ten sposób sprawdzić, że automat jest minimalny. Przykład: minimalny automat dla $(ab)^*$ ma 3 stany, są też 3 klasy abstrakcji: ϵ , a , aa . Uwaga: tu ważne jest, że definiujemy automat tak, że jest funkcja przejścia jest funkcją całkowitą.

Co więcej można pokazać następujący fakt. **Pytanie:** czy wszystkie automaty o minimalnej liczbie stanów są takie same?

Fakt

Wszystkie automaty deterministyczne dla L o $\text{indeks}(\sim_L)$ stanach są izomorficzne.

Izomorficzne, czyli można tak zmienić ich nazwy, żeby wyglądały tak samo. Tu formalnie zdefiniować izomorfizm (bijekcja t , że ...). Dlaczego tak jest? Niech A_L to będzie automat minimalny rozpoznający L , taki jak zdefiniowany w dowodzie tw. Myhill-Neroda. Niech A to jakiś inny automat rozpoznający L , taki, że jego liczba stanów jest równa liczbie stanów A_L . Chcemy pokazać, że A jest izomorficzny z A_L . Zauważmy, że wszystkie słowa, które w automacie A idą do stanu p muszą być z jednej klasy abstrakcji relacji $\sim_L$. Ale stanów jest $|Q_L|$, czyli do każdego stanu idą wszystkie słowa z pewnej klasy abstrakcji. No to już widać, że bijekcja, która przeprowadza p na p' takie, że jeśli $q_0 \xrightarrow{u} p$, to $p' = [u]$ jest izomorfizmem.

Pokazaliśmy więc, że automat minimalny jest w pewnym sensie kanoniczny. Okaże się, że możemy powiedzieć o nim więcej. Można go uzyskać z dowolnego automatu deterministycznego dla L poprzez "sklejanie" stanów. Zdefiniujemy teraz formalnie czym jest sklejanie stanów i jak to zrobić algorytmicznie.

Kongruencja

W zasadzie zawsze można skleić pewne zbiory stanów automatu i coś z tego wyjdzie. Przy czym przy niektórych sklejeniach automat może rozpoznawać zupełnie inny język. Przykład: skleamy stan początkowy ze śmietnikiem w automacie dla $(ab)^*$. Może też wyjść automat niedeterministyczny. Poza tym jak skleamy stan akceptujący z nieakceptującym, to jaki ma być stan sklejonny? Podamy teraz warunki na sensowne sklejenie, takie będzie nazywane kongruencją. Kongruencja to relacja równoważności $R \subseteq Q \times Q$ taka, że:

- 1) jeśli p i p' są w relacji R , to albo oba są akceptujące albo żaden
- 2) jeśli p i p' są w relacji R , $p \xrightarrow{a} q$, $p' \xrightarrow{a} q'$, to również q i q' są w relacji R

Sklejenie oczywiście skleja każdą klasę abstrakcji R w jeden stan.

Automat ilorazowy $A_R = (Q_R, \Sigma^*, q_R, F_R, \delta_R)$ jest zdefiniowany następująco:

- stany Q_R to klasy abstrakcji relacji R , czyli $\{[q]_R \mid q \in Q\}$
- stan początkowy q_R to klasa abstrakcji stanu początkowego q
- stany końcowe F_R to te klasy, które zawierają stany końcowe Q
- funkcja przejścia zawiera $p \xrightarrow{a} q$ jeśli przed sklejeniem z pewnym stanem była taka tranzycja (co więcej to znaczy, że w każdym stanie była taka tranzycja)

Łatwo pokazać, że $L(A_R) = L(A)$. Bieg przed sklejeniem przenosi się na bieg po sklejeniu, a bieg po sklejeniu przenosi się na pewien bieg przed sklejeniem (tu istotnie korzystamy z własności 1) i 2)). Zauważmy też, że A_R jest deterministyczny. Czyli fajnie, dzielenie przez kongruencję (czyli sklejenie klas abstrakcji tej kongruencji) nie zmienia języka i pozostawia automat deterministyczny. Teraz pokażemy, że dla każdego automatu istnieje kongruencja, że jak przez nią podzielimy, to dostaniemy automat minimalny. A ta relacja jest zdefiniowana w naturalny sposób. **Pytanie:** może ktoś ma pomysł jak? Sklejamy stany p i q wtedy, gdy $L(p, A) = L(q, A)$. Definiujemy: $L(p, A)$ to jest zbiór słów, które zostaną zaakceptowane w automacie A jeśli ustalimy p na stan początkowy. Zauważmy, że to jest faktycznie kongruencja (łatwo to zobaczyć). Załóżmy dla uproszczenia, że wszystkie stany automatu A są osiągalne (jeśli nie, to trzeba najpierw wywalić te nieosiągalne). Niech n to wielkość automatu A_L . Pomyślmy ile może być stanów po sklejeniu. Nie może być mniej niż n , bo to by dało mniejszy automat dla L . Ale nie może być też więcej niż n . Bo przypuśćmy, że do tych stanów prowadzą słowa $w_1, \dots, w_{n+1}$. Dla każdego w_i oraz w_j mamy stany q_i oraz q_j , do których prowadzą takie, że $L(p_i, A) \neq L(p_j, A)$. A więc $w_i \not\sim_L w_j$. A więc to dałoby przynajmniej $n+1$ klas abstrakcji relacji $\sim_L$, sprzeczność. A więc grup jest dokładnie n , czyli automat A_R ma dokładnie n stanów, czyli jest minimalny.

Podsumowanie

Czyli jest jeden kanoniczny automat minimalny, o najmniejszej liczbie stanów. Można go otrzymać z każdego innego poprzez sklejenie stanów, takie, że skleamy wszystkie stany p o równym języku $L(p, A)$.

Algorytm

Pytanie jak obliczyć ten minimalny automat. Trzeba po prostu umieć obliczyć dla których par stanów p oraz q zachodzi $L(p, A) = L(q, A)$. Powiemy, że stany p i q są n -rozdzielalne jeśli istnieje słowo w w długości co najwyżej n takie, że $w \in L(p, A)$, ale $w \notin L(q, A)$, albo odwrotnie. **Pytanie:** jeśli $L(p, A) \neq L(q, A)$, to jak długich słów potrzeba, żeby je rozróżnić?

Można pokazać, że n^2 wystarczy (pompowanie), ale zaraz okaże się, że nawet n wystarczy. Najpierw liczymy które stany są 0-rozróżnialne (akceptujące od nieakceptujących), a potem liczymy to dla coraz dłuższych słów. Jeśli mamy powiedziane, które pary stanów są k -rozróżnialne, to łatwo obliczyć które pary są $(k+1)$ -rozróżnialne, po prostu w jednym kroku musimy po tej samej literze osiągnąć parę stanów k -rozróżnialnych. Jak osiągniemy punkt stały, to się zatrzymujemy. Z tego co przed chwilą powiedzieliśmy mamy, że jeśli $L(p, A) \neq L(q, A)$, to są one n^2 -rozróżnialne. A więc nasz algorytm jest poprawny. A więc są one też tak naprawdę n -rozróżnialne (dygresja), bo punkt stały zostanie osiągnięty po co najwyżej n krokach. Czyli taka minimalizacja działa w $O(n)$ fazach po $O(n)$ czasu, czyli $O(n^2)$ czasu. Istnieje też lepszy algorytm (Hopcrofta), działający w czasie $O(n \log n)$. Do tej pory otwarty jest problem, czy da się to zrobić lepiej niż $O(n \log n)$, np. w czasie liniowym.

Przykład 1: automat dla $(ab)^*$ o 5 stanach. Zrobić dokładnie.

Przykład 2: automat liczący mod 3 o 9 stanach.

Wykład 5

Minimalizacja NFA

No dobrze, zminimalizowaliśmy automaty deterministyczne, ale co z minimalizacją automatów niedeterministycznych? Czy też tak się da? Niestety nie, nie ma nawet kanonicznego automatu minimalnego. Przykład: 2 automaty akceptujące język $K = \{ab, ac, ba, bc, ca, cb\}$. Możemy zrobić 2 automaty wielkości 5. Trzeba by w zasadzie pokazać, że nie da się tego zrobić automatem wielkości 4. Pokażemy to za chwilę.

Fooling set

Normalnie się tego nie robi, ale wg mnie to warto zrobić, bo jest fajne.

Tu zrobimy dygresję, jak pokazywać, że nie ma mniejszego automatu. **Pytanie:** dla deterministycznych jak pokazać, że nie ma automatu wielkości 7? Wskazać 8 słów nie w relacji M-N. Ale dla niedeterministycznego to nie wystarczy. Dlaczego? Wystarczy wziąć NFA mniejszy niż najmniejszy deterministyczny. Da się tak, np. n -ta litera od końca to a . NFA ma $n+1$ stanów. A DFA musi mieć przynajmniej 2^n , bo każde ze słów n literowych ma gdzieś różną literę, więc są nie w relacji M-N. Dobrze, to jak pokazywać, że NFA musi mieć przynajmniej ileś stanów? Z pomocą przychodzi następujący fakt.

Fakt (fooling set technique)

Niech L regularny. Jeśli istnieje n par słów $u_1, \dots, u_n$ oraz $v_1, \dots, v_n$ takich, że:

- $u_i v_i \in L$ dla każdego i
- $u_i v_k \notin L$ albo $u_k v_i \notin L$ dla każdych i, k ,

to wtedy minimalny automat niedeterministyczny dla L musi mieć przynajmniej n stanów.

Dowód

Ustalmy jakiś automat niedeterministyczny A dla języka L . Jeśli $u_i v_i \in L$, to znaczy, że musi być jakiś niepusty zbiór stanów S_i taki, że do każdego z tych stanów da się dojść po u_i ze stanu początkowego i z każdego z nich da się dojść po v_i do jakiegoś stanu końcowego. Z drugiego warunku wynika, że $S_i \cap S_j = \emptyset$. A zatem suma po S_i musi mieć przynajmniej n stanów.

c.n.d.

Zobaczmy teraz, że faktycznie nie ma dwóch automatów wielkości 4 dla $K = \{ab, ac, ba, bc, ca, cb\}$. Niech $u_1 = \epsilon$, $u_2 = a$, $u_3 = b$, $u_4 = c$, $u_5 = ab$ oraz $v_1 = ab$, $v_2 = b$, $v_3 = c$, $v_4 = a$, $v_5 = \epsilon$. Mamy $u_i v_i \in L$, a dla każdych dwóch różnych indeksów i, k któreś ze słów albo jest za krótkie albo ma dwie takie same litery.

Czyli faktycznie nie ma kanonicznego najmniejszego NFA.

Problemy decyzyjne

Jest kilka naturalnych pytań, które możemy zadać odnośnie danego języka regularnego.

Oto lista tych, którym się przyjrzymy:

- 1) dane słowo w , automat A , czy w należy do $L(A)$?
- 2) dany automat A , czyli $L(A)$ pusty?
- 3) dane automaty A, B , czy $L(A) = L(B)$?
- 4) dane automaty A, B , czy $L(A) \subseteq L(B)$?

Każde z tych pytań można zadać dla automatu deterministycznego i niedeterministycznego.

Przyjrzyjmy się tym problemom. Jak zrobić 1) dla DFA? Po prostu lecimy po słowie, wychodzi w $O(|w| |A|)$. **Pytanie:** a jak dla NFA? Też lecimy po słowie, tylko trzymamy zbiór stanów do którego możemy dojść. Też PTIME, dokładniej $O(|w| |A|^2)$.

Problem 2) też łatwo w PTIME dla obu, patrzymy czy jest ścieżka z jakiegoś początkowego do jakiegoś końcowego. Jak dla 3)? Dla DFA łatwo, dopełniamy drugi, przecinamy i badamy niepustość. A dla NFA? To się okaże trudne, nie znamy algorytmu wielomianowego. Podobnie zresztą sprawdzanie, czy $L(A) = \Sigma^*$, nie znamy tu algorytmu wielomianowego i podejrzewamy, że go nie ma (problem jest PSPACE-trudny, pod koniec wykładu być może powiemy co to znaczy). Można zdeterminizować i zrobić to, co dla DFA, działa w czasie wykładniczym. Dla 4) robimy podobnie jak dla 3), bo to jest równoważne pytaniu, czy $L(A) \cap \bar{L(B)} \neq \emptyset$.

Algebraiczne podejście dla języków regularnych

Teraz zrobimy temat nieco dodatkowy, tzn. zajmiemy się algebraicznym podejściem do języków regularnych. Tego się zwykle nie robi, ale my w tym roku zrobimy. Rozważmy pewien automat deterministyczny A . Bardzo przydatnym pojęciem dla słowa $w \in \Sigma^*$ jest funkcja $f_w: Q \rightarrow Q$, która mówi które stany przechodzą na które po przejściu przez słowo w . To znaczy definiujemy ją tak, że $f_w(p) = q$ wtw. $p \xrightarrow{w} q$ w automacie A . Dużo zadań można rozwiązać używając tej funkcji, a co ważniejsze można łatwiej myśleć używając tego pojęcia.

Przykład:

Niech L regularny, pokaż, że $\sqrt{L} = \{w \mid ww \in L\}$ jest regularny.

Weźmy DFA A dla języka L . Niech A' oblicza funkcję po słowie w , jego stany to różne możliwe funkcje. $f_{\epsilon} = \text{id}$, a potem łatwo liczyć. Stany akceptujące to te funkcje f , dla których $f(f(q_0)) \in F$. Tak samo można zrobić dla $\sqrt[7]{L}$ i sporo innych podobnie.

Zobaczmy, że takie funkcje są monoidem, czyli strukturą $(M, *)$ taką, że:

- istnieje element neutralny
- każde dwa elementy można przemnożyć i dostać znów element z dziedziny

Grupy to szczególny rodzaj monoidów, one mają jeszcze elementy odwrotne, monoidy niekoniecznie. Okazuje się, że takie spojrzenie jest bardzo przydatne. Pokażemy teraz twierdzenie, które będzie w nowy sposób charakteryzowało języki regularne. Potem wprowadzimy nowy rodzaj automatów - automaty dwukierunkowe i używając monoidów pokażemy, że one nie umieją zrobić więcej niż automaty jednokierunkowe.

Homomorfizm między dwoma monoidami $(M, *1)$ i $(N, *2)$ to funkcja $h: M \rightarrow N$ taka, że dla każdych $m, m' \in M$ zachodzi: $h(m1) *2 h(m2) = h(m1 *1 m2)$.

Powiemy, że język L jest rozpoznawany przez monoid M jeśli istnieje homomorfizm $h: \Sigma^* \rightarrow M$ oraz podzbiór akceptujący $F \subseteq M$ taki, że

$$h^{-1}(F) = L,$$

czyli słowa są akceptowane jeśli przechodzą na zbiór akceptujący.

Przykłady jak można rozpoznać języki regularne przez monoid:

1. Słowa długości przystającej do 3 modulo 1. Monoid: $\{0, 1, 2\}$ z działaniem dodawania modulo 3. Funkcja $h: \Sigma^* \rightarrow \{0, 1, 2\}$, zwraca długość modulo 3. To jest homomorfizm, bo $h(uv) = h(u) + h(v)$. Dajemy $F = \{1\}$. Mamy wtedy $L = h^{-1}(F)$.
2. Słowa postaci $a^* b^* c^*$. Monoid: $\{\text{nie}, a, b, c, ab, bc, \text{zle}\}$. Działanie na logikę. $F = M \setminus \{\text{zle}\}$. $L = h^{-1}(F)$.
3. Słowa o podciągu abc . Monoid: $\{\text{jest, nie ma}\} \times \{\text{pref. bc, pref c, pref nie c}\} \times \{\text{suf. ab, suf. a, suf. nie a}\}$. $F = \{\text{jest}\} \times \text{all} \times \text{all}$. Działa.
4. Słowa $a^n b^n$. Monoid: $\{a \text{ długości } i, b \text{ długości } i, a^i b^j, b \text{ przed } a\}$. Ale tu nieskończony.
5. Dowolny język. Monoid: $(\Sigma^*, \text{konkatenacja})$, $F = L$, $h = \text{id}$. Wtedy $L = h^{-1}(F)$. Ale tu nieskończony.

Tw.

L jest regularny wtw. L jest rozpoznawalny przez pewien monoid **skończony**.

Dowód

Najpierw w prawo. Niech L regularny i A to automat deterministyczny, który go rozpoznaje, ze zbiorem stanów Q . Wtedy robimy $M = (Q^*Q, \text{złożenie})$, a F to funkcje t., że stan początkowy A

przeprowadzają na jakiś akceptujący. Mamy też, że $h(w) = f_w$ jest homomorfizmem, bo $h(u)h(v) = h(uv)$. Poza tym $h^{-1}(F) = L$, co kończy dowód tej implikacji.

Teraz w lewo. Przypuśćmy, że L jest rozpoznawalny przez monoid $(M, *)$. Automat A po prostu liczy do jakiego elementu monoidu dane słowo się wyliczy. Czyli zbiór stanów to M . Stan początkowy to $h(\epsilon)$. Gdy mamy $h(w)$, to po literze a przechodzimy do $h(w) \circ h(a)$. Stany akceptujące to te z F .
c.n.d.

Wykład 6

To jest kolejna charakteryzacja języków regularnych (już 4-ta), co pokazuje, że są rzeczywiście interesującym obiektem. Teraz zdefiniujemy automaty dwukierunkowe. One będą zachowywały się nieco inaczej niż normalne automaty. Z powodów technicznych stany będą nie między literami, ale na literach. Przykład automatu, który rozpoznaje $\Sigma^* a \Sigma^{n-1}$, czyli n -ta litera od końca to a . Idea: on idzie do końca, a potem się wraca licząc do n i sprawdza, czy tam jest a . Może to nawet zrobić deterministycznie!

Formalnie $A = (\Sigma, |-, -|, Q, I, F, \delta)$, gdzie

- Σ to alfabet, $|-$ oraz $-|$ to specjalne symbole ograniczające słowo odpowiednio z lewej i z prawej strony
- Q to zbiór stanów
- I to stany początkowe
- F to stany akceptujące
- δ to relacja przejścia, mamy $\delta \subseteq Q \times (\Sigma \cup \{|-, -|\}) \times Q \times \{-1, 0, 1\}$.

Interpretacja relacji przejścia jest jak zwykle, przy czym $\{-1, 0, 1\}$ oznacza w którą stronę automat teraz pójdzie. Zaczynamy zawsze w stanie początkowym na pierwszej literze. Kończymy w stanie końcowym.

Przykład:

1. Napiszmy automat dla języka: n -ta litera od końca to a .
 $Q = \{p, q_1, \dots, q_n, akc, rej\}$
 $A = (\Sigma, |-, -|, Q, p, \{akc\}, \delta)$
 Do δ należą:
 $(p, \text{wszystko oprócz } -|, p, 1)$ - idziemy w prawo
 $(p, -|, q_1, -1)$
 $(q_i, \text{wszystko}, q_{i+1}, -1)$
 $(q_n, a, akc, 0)$
 $(q_n, \text{nie } a, rej, 0)$
 $(rej, \text{wszystko}, rej, 0)$
2. Każdy język regularny da się zrobić. Po prostu stan, który powinien być przed literą stawiamy na literze. Stan końcowy odczytujemy na literze $-|$.

Teraz pytanie brzmi, czy da się zrobić coś więcej niż języki regularne? Nie, nie da się, już wcześniej mówiłem. Można to łatwo pokazać używając monoidów.

Tw.

Automaty dwukierunkowe rozpoznają dokładnie języki regularne.

Dowód

To, że akceptują języki regularne, to już pokazaliśmy powyżej.

Teraz rozważmy pewien automat dwukierunkowy A . Najpierw założmy dla uproszczenia, że A jest deterministyczny, czyli, że ma dokładnie jedno przejście dla ustalonego stanu i litery. Pokażemy, że jego język da się rozpoznać pewnym monoidem skończonym, co zakończy dowód. Zawsze jak projektujemy taki monoid musimy się zastanowić co właściwie chcemy wiedzieć o słowie, czego nam potrzeba. Wtedy monoid będzie zbiorem wszystkich możliwych takich informacji, a słowo będzie mapowane przez homomorfizm na odpowiednią wartość. Tutaj żeby wiedzieć wszystko o słowie wystarczy wiedzieć co się stanie jeśli automat wejdzie do niego z lewej (lub z prawej) w danym stanie - czy wyjdzie z lewej lub z prawej, zaakceptuje lub się zapętli. A więc dla słowa w zdefiniujemy funkcję, która mówi, że jeśli zaczynamy w słowie na pierwszej pozycji z lewej/prawej w danym stanie, to jakie będzie możliwe wyniki: wyjście z lewej/prawej w danym stanie, akceptacja, zapętlenie się w środku. Czyli $f_w: Q \times \{L, P\} \rightarrow Q \times \{L, P\} \cup \{akc, pętla\}$. Takie funkcje można mnożyć tak, żeby wyszedł homomorfizm, bo da się odtworzyć co funkcja robi na słowie uv , gdy wiemy co robi na u i co robi na v . A więc monoid to (funkcje, operacja na nich). Funkcje akceptujące to te, dla których jak zaczynamy na pierwszej pozycji z lewej w stanie początkowym, to zaakceptujemy. A więc pokazaliśmy, że deterministyczne automaty dwukierunkowe akceptują tylko języki regularne. Analogicznie można zrobić dla niedeterministycznych, tylko, że trzeba pamiętać relację, czyli wszystkie możliwe przejścia. Wtedy zbiór akceptujący to taki, że istnieje zaczęcie ze zbioru stanów początkowych takie, że to doprowadzi do akceptacji.

c.n.d.

Problemy otwarte dotyczące automatów dwukierunkowych

Jak duży jest wzrost gdy przekształcamy automat niedeterministyczny (1 lub 2 kierunkowy) na automat deterministyczny 2 kierunkowy?

Gramatyki bezkontekstowe

Zaczynamy teraz drugi duży dział na tym przedmiocie, dotyczący języków bezkontekstowych. Mogą być one opisywane albo przez gramatyki bezkontekstowe albo przez automaty ze stosem. Zaczniemy od gramatyk bezkontekstowych, bo one są prostsze do zrozumienia.

Języki bezkontekstowe to większa klasa niż regularne. Okaże się, że każdy język regularny jest bezkontekstowy, ale pewne języki bezkontekstowe nie są regularne.

Zacznijmy od przykładu, palindromy nad dowolnym skończonym alfabetem (np. $\{a, b\}$) są językiem bezkontekstowym. Gramatyka dla nich to:

$P \rightarrow aPa$

$P \rightarrow bPb$

$P \rightarrow a$

$P \rightarrow b$

$P \rightarrow \epsilon$

Idea jest taka, że zaczynamy od symbolu P , a potem aplikujemy reguły. Język to zbiór rzeczy, które możemy dostać. Np. możemy uzyskać słowo $abbba$ w następujący sposób:

$P \rightarrow aPa \rightarrow abPba \rightarrow abbba$

Definicja

Teraz formalna definicja. Gramatyka (T, N, S, δ) składa się ze:

- zbioru terminali T (symboli końcowych, czyli literek)
- zbioru nieterminali N (symboli pomocniczych, czasem mówimy zmiennych, dużych liter)
- symbolu startowego $S \in N$
- zbioru produkcji $\delta \subseteq N \times (N \cup T)^*$

Powiedzieć co jest czym w gramatyce powyżej i że tam był tylko jeden nieterminal, ale w ogólności może być więcej.

Produkcje (czasem mówimy reguły, albo też może inaczej) zapisujemy często jako $X \rightarrow \alpha$ zamiast $(X, \alpha) \in \delta$. Reguły rozszerzamy też do relacji na ciągach nieterminali, tzn. np. gdy $X \rightarrow \alpha$, to piszemy również $X\gamma \rightarrow \beta\alpha\gamma$.

Wyprowadzenie (derywacja) to ciąg

$\alpha_0 \rightarrow \dots \rightarrow \alpha_n$

taki, że $\alpha_i \rightarrow \alpha_{i+1}$, $\alpha_0 = S$ oraz $\alpha_n \in T^*$.

Wówczas mówimy, że α_n da się wyprowadzić w gramatyce G .

Język $L(G)$ gramatyki G to zbiór wszystkich słów, które dadzą się wyprowadzić w gramatyce G .

Czasem wyprowadzeniem nazywamy też wyprowadzenie bez założeń $\alpha_0 = S$ i $\alpha_n \in T^*$, nie będziemy wprowadzać dodatkowych definicji.

Przykłady

- $a^n b^n$: $P \rightarrow aPb \mid \epsilon$
- $\{w \mid \#a(w) = \#b(w)\}$: $P \rightarrow aPb \mid bPa \mid PP \mid \epsilon$
(tu dowód potrzebny - robimy wykres $\#a(w) - \#b(w)$, patrzymy gdzie przecina 0)
- wyrażenia nawiasowe $P \rightarrow PP \mid (P) \mid \epsilon$
- $a^n b^k$, $n \neq k$ $P \rightarrow A \mid B$, $A \rightarrow aA \mid aR$, $B \rightarrow Bb \mid Rb$, $R \rightarrow aRb \mid \epsilon$
- wyrażenia arytmetyczne $W \rightarrow W + W \mid W * W \mid (W) \mid 0 \mid 1$

- palindromy długości podzielnej przez 3

$$X_0 \rightarrow a X_1 a \mid b X_1 b \mid \text{eps}, X_1 \rightarrow a X_2 a \mid b X_2 b \mid a \mid b, \\ X_2 \rightarrow a X_0 a \mid b X_0 b$$

Bezkontekstowe vs regularne

Fakt

Każdy język regularny jest bezkontekstowy

Dowód

Najpierw przykład. Weźmy automat rozpoznający $(ab)^*c(b+c)^*$. Ma 3 stany, nie licząc śmietnika, bo go tu nie potrzebujemy: q_0 , q_1 , q_2 , stan q_0 początkowy, stan q_2 akceptujący. Ma tranzyccje $q_0 \xrightarrow{a} q_1$, $q_1 \xrightarrow{b} q_0$, $q_0 \xrightarrow{c} q_2$, $q_2 \xrightarrow{b} q_2$, $q_2 \xrightarrow{c} q_2$.

Robimy gramatykę. Terminale $\{a, b, c\}$, nieterminale: X_0, X_1, X_2 . Początkowy nieterminał X_0 .

Produkcje:

$$X_0 \rightarrow a X_1$$

$$X_1 \rightarrow b X_0$$

$$X_0 \rightarrow c X_2$$

$$X_2 \rightarrow b X_2$$

$$X_2 \rightarrow c X_2$$

$$X_2 \rightarrow \text{eps}$$

Te wcześniejsze odpowiadają tranzyccjom automatu, ta późniejsza faktowi, że q_2 jest akceptujący.

Zobaczmy jak wyprowadzić $abcb$.

$$X_0 \rightarrow a X_1 \rightarrow a b X_0 \rightarrow a b c X_2 \rightarrow a b c b X_2 \rightarrow a b c b b X_2 \rightarrow a b c b b$$

Widać, że to odpowiada biegowi automatu.

W ogólności dla automatu $(Q, \Sigma, q_0, F, \delta)$ robimy

terminale = Σ , nieterminale = $X_q, q \in Q$ (albo Q , ale to by się myliło), nieterminał

początkowy to X_{q_0} , dla każdej tranzyccji $p \xrightarrow{a} q$ robimy produkcję $X_p \rightarrow a X_q$ oraz

dla każdego stanu końcowego p robimy produkcję $X_p \rightarrow \text{eps}$.

c.n.d.

Drzewa wyprowadzeń

Przykłady:

- dla gramatyki $X \rightarrow a X b \mid b X a \mid X X \mid \text{eps}$ oraz słowa $aababbb$
- dla gramatyki palindromów długości podzielnej przez 3 oraz słowa $abbaaabb$

Pytania:

- czy dla każdej derywacji istnieje odpowiadające drzewo derywacji? (TAK)
- czy dokładnie jedno? (TAK)
- czy dla każdego drzewa derywacji istnieje odpowiadająca derywacja? (TAK)
- czy dokładnie jedna? (NIE)

Ale można dla każdego drzewa derywacji wyszczególnić jedną derywację, tzw. lewostronną derywację. To odpowiada stosowaniu zawsze reguły dla najbardziej lewego nieterminala.

Wykład 7

Lemat o pompowaniu

Podobnie jak przy lemacie o pompowaniu dla języków regularnych łatwiej jest tu zapamiętać dowód niż sformułowanie. Ja np. nie pamiętam sformułowania, tylko wyprowadzam je pamiętając ideę dowodu.

Dla każdego języka bezkontekstowego L istnieje taka stała N , że dla każdego słowa $w \in L$ długości co najmniej N istnieje podział $w = xsyztz$ taki, że:

- $st \neq \epsilon$
- $|sy| \leq N$
- dla każdego i zachodzi $x s^i y t^i z \in L$

Dowód

Ustalmy gramatykę bezkontekstową G taką, że $L = L(G)$. Niech n to liczba nieterminali w gramatyce i niech k to maksymalna liczba symboli po prawej stronie produkcji. Zauważmy, że jeśli mamy drzewo wyprowadzenia o głębokości nie większej niż n , to słowo z niego wyprowadzone ma długość co najwyżej n^k . Ustalmy $N = n^{k+1} + 1$. Niech $w \in L$ takie, że $|w| \geq N$. Wówczas drzewo wyprowadzenia dla w na głębokość co najmniej $n+2$, czyli jest pewna ścieżka, która ma długość co najmniej $n+2$. Na tej ścieżce jest więc co najmniej $n+1$ nieterminali, czyli któryś powtarza się co najmniej 2 razy. Niech będzie to X , pod nim jest drugi X . A więc mamy wyprowadzenie $X \rightarrow s X t$ na $st \neq \epsilon$. Tu narysować ogólną sytuację. Jak wstawimy (lub usuniemy) takie drzewo wyprowadzenia do tego co mamy to dostaniemy $x s^i y t^i z \in L$. Dodatkowo możemy wziąć najniższą taką parę, wtedy dostaniemy $|sy| \leq n^{k+1}$, czyli $|sy| \leq N$.
c.n.d.

Ten lemat o pompowaniu, podobnie jak lemat o pompowaniu dla języków regularnych nie zawsze działa. Z tego co wiem, to nie żadnego warunku, który zawsze działa. Ale są lepsze lematy niż lemat o pompowaniu, działające na podobnej zasadzie. Zrobimy taki jeden, lemat Ogdena. Tam po prostu pompuje się uważniej.

Lemat Ogdena

Dla każdego języka bezkontekstowego L istnieje taka stała N , że dla każdego słowa $w \in L$ z co najmniej N zaznaczonymi pozycjami istnieje podział $w = xsyztz$ taki, że:

- st na przynajmniej jedną zaznaczoną pozycję
- syt ma mniej niż N zaznaczonych pozycji
- dla każdego i zachodzi $x s^i y t^i z \in L$

Dowód jest w zasadzie taki sam. Patrzymy na ścieżki idące od zaznaczonych pozycji do korzenia i drzewo stworzone przez te ścieżki, wierzchołki drzewa dajemy tam, gdzie schodzą się ścieżki. To drzewo ma stopień co najwyżej k , więc dla odpowiednio dużego N , tak jak dla zwykłego lematu o pompowaniu, istnieje pewna ścieżka na której powtarza się symbol. Tam pompujemy, tak jak wcześniej.

Zauważmy, że dla wszystkich pozycji zaznaczonych lemat Ogdena jest równoważny lematowi o pompowaniu.

Przykłady

1. Rozważmy język $L = \{a^n b^n c^n \mid n \in \mathbb{N}\}$. Weźmy stałą K z lematu. Rozważmy $a^K b^K c^K$. Ani s ani t nie może zachodzić na obszary z różnych liter. Zatem s i t zawierają takie same litery. Zatem $x s^2 y t^2 z$ będzie zawierać więcej niektórych liter (bo zwiększymy liczbę tylko co najwyżej dwóch liter).
2. Rozważmy język $L = \{a^i b^j c^k d^l \mid j = k = l \text{ lub } i = 0\}$. Lemat o pompowaniu nie działa, bo zawsze może się okazać, że pompujemy litery a - i wtedy jest to nieszkodliwe, słowo nadal jest w języku. Ale z lematu Ogdena da się to zrobić. Zaznaczamy wszystko oprócz liter a i wtedy jest ok, tak samo jak w przykładzie pierwszym robimy.

Jednoznaczność

Pytanie na ile wyprowadzenia są unikalne dla słowa.

Można to pytania zadać na różnych poziomach:

- jedno wyprowadzenie dla każdego słowa (łatwo pokazać, że nie zawsze jest)
 $X \rightarrow XX \rightarrow abX \rightarrow abab, X \rightarrow XX \rightarrow Xab \rightarrow abab$
- no ale te wyprowadzenia odpowiadają jednemu drzewu derywacji, może prawdą jest, że zawsze jest jedno drzewo derywacji w danej gramatyce (łatwo pokazać, że nie zawsze jest)
 $X \rightarrow aX \mid Xa \mid \epsilon, X \rightarrow aX \rightarrow a$ oraz $X \rightarrow Xa \rightarrow a$
- no dobrze, ale widać, że tu jest inna gramatyka, która byłaby jednoznaczna: $X \rightarrow aX \mid \epsilon$. Dla każdego języka bezkontekstowego jest gramatyka, która nie jest jednoznaczna, wystarczy z nieterminała X zrobić dwie wersje: X_1 oraz X_2 . Ale ciekawym pytaniem jest: czy dla każdego języka bezkontekstowego jest pewna gramatyka, że tam jest dla każdego słowa jedno drzewo derywacji. **Pytanie:** czy zawsze istnieje taka gramatyka? Odpowiedź brzmi: NIE. Przykład języka to $\{a^i b^n c^n\} \cup \{a^n b^n c^i\}$. Słowo $a^n b^n c^n$ powinno mieć zawsze dwie derywacje. Bez dowodu.

Postacie normalne

Okazuje się, że każdą gramatykę można sprowadzić do postaci, które są prostsze w obsłudze. Najbardziej znaną taką postacią jest postać normalna Chomsky'ego.

Gramatyka bezkontekstowa jest w postaci normalnej Chomsky'ego gdy występują tam tylko reguły następującej postaci:

$X \rightarrow YZ$, gdzie Y i Z to nieterminale

$X \rightarrow a$, gdzie a to terminal

$S \rightarrow \text{eps}$

Fakt

Dla każdego języka bezkontekstowego L istnieje gramatyka bezkontekstowa w postaci Chomsky'ego G taka, że $L(G) = L$.

Dowód

Jeśli $\text{eps} \notin L$, to dodajemy regułę $S \rightarrow \text{eps}$. Będziemy teraz starali się zdefiniować gramatykę dla języka $L \setminus \{\text{eps}\}$, czyli zakładamy, że eps nie należy do L .

Będziemy po kolei przekształcać naszą pierwotną gramatykę, tak, by na końcu zostały jedynie reguły postaci $X \rightarrow YZ$, $X \rightarrow a$ oraz $X \rightarrow \text{eps}$.

Krok 1. Najpierw dla każdego terminala a wprowadzamy specjalny nieterminał A i dodajemy regułę $A \rightarrow a$. Następnie każde wystąpienie litery a po prawej stronie (oprócz tej właśnie reguły) zastępujemy przez A . Po tym ruchu wszystkie reguły są albo postaci $X \rightarrow \text{eps}$, $X \rightarrow a$ albo postaci $X \rightarrow Y_1 \dots Y_k$. Problem polega na tym, że wciąż może być $k = 1$ lub $k > 2$.

Krok 2. Eliminacja reguł postaci $X \rightarrow \text{eps}$.

Patrzmy z których zmiennych możemy wyprodukować słowo puste, niech będzie to zbiór P . Następnie dla każdej reguły $X \rightarrow \alpha$ dodajemy reguły postaci $X \rightarrow \beta$ dla każdego niepustego β , które powstało z α poprzez niektórych zmiennych ze zbioru P . Następnie kasujemy wszystkie reguły $X \rightarrow \text{eps}$. Nie zmieniliśmy przez to języka, to po prostu robimy taką derywację, że nigdy nie produkujemy zmiennych, które miałyby potem wyprodukować słowo puste.

Krok 3. Eliminacja reguł postaci $X \rightarrow Y$.

Patrzmy które zmienne mogą przejść na które ciągiem takich kroków.

Jeśli mamy $Y_1 \rightarrow Y_2 \rightarrow \dots \rightarrow Y_k$, to mówimy, że Y_1 może przejść na Y_k .

Dla każdej pary (X, Y) takiej, że X może przejść na Y oraz dla każdej reguły $Y \rightarrow \alpha$ dodajemy również regułę $X \rightarrow \alpha$. Widać, że to nie zwiększa mocy. Następnie eliminujemy wszystkie przejścia $X \rightarrow Y$. To nie zmienia języka, to dla każdej derywacji w starej formie, która używałaby $Y_1 \rightarrow \dots \rightarrow Y_k \rightarrow \alpha$ robimy od razu $Y_1 \rightarrow \alpha$.

Krok 3. Eliminacja reguł postaci $X \rightarrow Y_1 \dots Y_k$ dla $k > 2$.

Dla każdej reguły $X \rightarrow Y_1 \dots Y_k$ dla $k > 2$ robimy nowe zmienne $Z_1, \dots, Z_{(k-2)}$ i regułę $X \rightarrow Y_1 Z_1$, dla każdego $i < k-2$ regułę $Z_i \rightarrow Y_{(i+1)} Z_{(i+1)}$ oraz regułę $Z_{(k-2)} \rightarrow Y_{(k-1)} Y_k$.

Potem usuwamy oryginalną regułę $X \rightarrow Y_1 \dots Y_k$.

c.n.d.

Postać normalna Greibach (bez dowodu)

Można też każdą gramatykę G taką, że $L(G)$ nie zawiera słowa pustego zamienić na gramatykę z regułami tylko postaci:

$X \rightarrow a Y_1 \dots Y_k$ dla $k \geq 0$.

Nie podajemy dowodu, jest trochę techniczny.

Automaty ze stosem

Najpierw przykład: dla palindromów

Tranzycje to:

$(q, S) \xrightarrow{-a} (q, AS)$

$(q, S) \xrightarrow{-b} (q, BS)$

$(q, X) \rightarrow (p, X)$, X dowolne

$(p, A) \xrightarrow{-a} (p, \text{eps})$

$(p, B) \xrightarrow{-b} (p, \text{eps})$

$(p, S) \rightarrow (f, S)$

Stan q jest początkowy, a f akceptujący. Startujemy z konfiguracji (q, S) .

Automat (niedeterministyczny) ze stosem $A = (\Sigma, Q, q_0, F, \Gamma, s_0, \delta)$ składa się z:

- skończonego alfabetu Σ
- zbioru stanów Q
- stanu początkowego $q_0 \in Q$
- zbioru stanów końcowych $F \subseteq Q$
- alfabetu stosowego Γ (zbioru symboli, które mogą pojawiać się na stosie)
- symbolu początkowego $s_0 \in \Gamma$
- zbioru tranzycji $\delta \subseteq Q \times \Gamma \times (\Sigma \cup \text{eps}) \times Q \times \Gamma^*$

Konfiguracja automatu ze stosem to stan $z \in Q$ oraz ciąg symboli stosowych $\alpha \in \Gamma^*$.

Automat startuje w konfiguracji (q, s_0) . Następnie zmienia stan i czubek stosu wedle tranzycji.

Kończy w konfiguracji (f, α) o ile $f \in F$, α może być dowolna.

Interpretacja tranzycji (q, A, a, q', α) jest następująca: w stanie q , o ile na czubku stosu jest symbol stosowy A , to możemy wczytać z wejścia literę a , zmienić stan na q' , a na stosie zamiast A umieścić ciąg symboli stosowych α . Będziemy tę tranzycję zapisywać jako

$(q, A) \xrightarrow{-a} (q', \alpha)$. Jeśli $c_0 \xrightarrow{-a_1} c_1 \dots c_{n-1} \xrightarrow{-a_n} c_n$, to piszemy $c_0 \xrightarrow{-w} c_n$ dla $w = a_1 \dots a_n$. Jeśli $(q, s_0) \xrightarrow{-w} (f, \alpha)$ dla pewnego $f \in F$, to mówimy, że A akceptuje w .

Język automatu, $L(A)$, to zbiór wszystkich słów w , które A akceptuje.

Wykład 8

Przykład

Robimy automat dla języka palindromów nad $\{a, b\}$ o długości podzielnej przez 3.

$\Sigma = \{a, b\}$

$Q = \{p_0, p_1, p_2, q_0, q_1, q_2\}$

początkowy - p_0

$F = \{q_0\}$

$\Gamma = \{P, A, B\}$

$s_0 = P$

Idea jest taka, że w stanach p_i będziemy w pierwszej połowie słowa, a w q_i w drugiej.

Dodatkowo indeks liczy modulo 3. Czubek stosu jest z lewej.

Tranzycje to:

$(p_i, X) \xrightarrow{-a} (p_j, AX)$ o ile $j = i+1 \bmod 3, X \in \{A, B, P\}$

$(p_i, X) \xrightarrow{-b} (p_j, BX)$ o ile $j = i+1 \bmod 3, X \in \{A, B, P\}$

$(p_i, X) \xrightarrow{-x} (q_j, X)$ o ile $j = i+1 \bmod 3, x \in \{a, b\}, X \in \{A, B, P\}$

$(p_i, X) \xrightarrow{-\epsilon} (q_i, X)$ o ile $X \in \{A, B, P\}$

$(q_i, A) \xrightarrow{-a} (q_j, \epsilon)$ o ile $j = i+1 \bmod 3$

$(q_i, B) \xrightarrow{-b} (q_j, \epsilon)$ o ile $j = i+1 \bmod 3$

$(q_0, P) \xrightarrow{-\epsilon} (f, \epsilon)$

Przykładowy bieg:

$(p_0, P) \xrightarrow{-a} (p_1, AP) \xrightarrow{-b} (p_2, BAP) \xrightarrow{-b} (p_0, BBAP) \rightarrow (q_0, BBAP) \xrightarrow{-b} (q_1, BAP) \xrightarrow{-b} (q_2, AP) \xrightarrow{-a} (q_0, P) \rightarrow (f, \epsilon)$

Zatem abbbbba należy do języka.

Zauważmy, że nasz automat ma kilka własności:

- każda tranzycja albo wrzuca jeden symbol na stos (push) albo zdejmuję ze stosu (pop),
- jest tylko jeden stan końcowy
- każde obliczenie akceptujące kończy się usunięciem symbolu początkowego i zaakceptowaniu pustym stosiem
- symbol początkowy s_0 nigdy nie jest wrzucany na stos, a zawsze jest usuwany tuż przed zaakceptowaniem

Jest to ogólniejszy fakt. Każdy język akceptowany przez pewien automat można zaakceptować takim właśnie automatem.

Warianty automatów ze stosiem

Tw.

Dla każdego automatu ze stosem A istnieje automat ze stosem A' taki, że $L(A) = L(A')$ oraz dodatkowo A' spełnia warunki:

- każda tranzycja jest albo typu pop albo typu push
- jest tylko jeden stan końcowy
- każde obliczenie akceptujące akceptuje pustym stosem
- symbol początkowy nigdy nie może być wrzucony na stos (nie ma takiej tranzycji) i zawsze jest usuwany tuż przed akceptacją

Uwaga: zauważmy, że w takim razie automat A' akceptuje wtw. gdy jest stos jest pusty. Można zatem zmienić jego warunek akceptacji mówiąc, że akceptuje, gdy stos jest pusty. Będziemy więc od tej pory mogli rozważać dwa alternatywne warunki akceptacji:

- przez stan
- przez pusty stos

Zobaczmy, że akceptacja przez pusty stos nie jest mocniejsza, bo możemy dla każdego automatu, który akceptuje przez pusty stos wrzucić coś na dół i jak to zobaczymy na dole to przejść do specjalnego stanu akceptującego (nowego).

Można by teoretycznie rozważać też warunek przez stan oraz pusty stos, ale to nie jest wcale wygodniejsze.

Dowód (twierdzenia)

Dodajemy nowy stan początkowy q_0 i symbol początkowy P_0 , robimy też tranzycję $(q_0, P_0) \rightarrow (q_0', P_0' P_0)$, gdzie q_0' to stary stan początkowy, a P_0' to stary symbol początkowy.

Robimy też nowe stany końcowy f i normalny q oraz tranzycje

$(f', X) \rightarrow (q, X)$ dla każdego starego stanu akceptującego f' i zmiennej X

$(q, X) \rightarrow (q, \text{eps})$ dla każdej zmiennej X różnej od P_0

$(q, P_0) \rightarrow (f, \text{eps})$

One pozwalają na opróżnienie stosu przed akceptacją.

Teraz pozostaje jedynie pokazać, że wystarczy ograniczyć się do tranzycji typu push i pop.

Weźmy tranzycję $(p, X) \xrightarrow{a} (q, Y_1 \dots Y_k)$. Jeśli $k = 0$, to jest typu pop, w innym wypadku będziemy ją przerabiać. Dodajemy nowe stany $r_1, \dots, r_k$ oraz tranzycje:

$(p, X) \xrightarrow{a} (r_1, \text{eps})$

$(r_1, Y) \rightarrow (r_2, Y_1 Y)$ dla dowolnego Y

$(r_i, Y_{\{i-1\}}) \rightarrow (r_{i+1}, Y_i Y_{\{i-1\}})$ dla $i < k$

$(r_k, Y_{\{k-1\}}) \rightarrow (q, Y_k Y_{\{k-1\}})$

Przerobiliśmy więc dowolną tranzycję na serię popów (jednego) i pushy. Stany $r_1, \dots, r_k$ muszą być różne dla różnych tranzycji.

c.n.d.

Równoważność automatów ze stosem i gramatyk

Pokażemy teraz, że automaty ze stosem i gramatyki akceptują te same języki.

Tw.

Język L jest rozpoznawany przez pewną gramatykę bezkontekstową wtw. jest akceptowany przez pewien automat ze stosem.

Dowód

Najpierw pokażemy implikację " $\Rightarrow$ ", która jest łatwiejsza. Niech L będzie rozpoznawany przez pewną gramatykę bezkontekstową G . Możemy bez straty ogólności założyć, że G jest w postaci normalnej Chomsky'ego. Czyli ma reguły postaci $S \rightarrow \text{eps}$ (potencjalnie), $X \rightarrow a$, $X \rightarrow YZ$. Robimy więc automat ze stosem A taki, że $L(A) = L$, będziemy zakładać, że akceptuje on przez pusty stos.

Stan jest jeden: $Q = \{q\}$. Jest on zarówno początkowy jak i akceptujący.

Symbol początkowy to S . Tranzycje są następujące. Po pierwsze mamy tranzycję $(q, S) \rightarrow (q, \text{eps})$, o ile taka jest w gramatyce. Po drugie dla każdej reguły $X \rightarrow a$ mamy tranzycję $(q, X) \xrightarrow{-a} (q, \text{eps})$.

Po trzecie dla każdej reguły $X \rightarrow YZ$ mamy tranzycję $(q, X) \rightarrow (q, YZ)$. Widać, że język jest taki sam.

Teraz pokażemy trudniejszą implikację " $\Leftarrow$ ". Niech L będzie rozpoznawany przez pewien automat ze stosem $(\Sigma, Q, q_0, F, \Gamma, \delta)$ akceptujący przez pusty stos.

Zaprojektujemy gramatykę, która będzie rozpoznawała ten sam język. Zbiór nieterminali to $\{X_{p,q}^s \mid p, q \in Q, s \in \Gamma\}$ oraz specjalny symbol S . Idea jest taka, że język $L(X_{p,q}^s)$ będzie równy $\{w \mid (p, s) \xrightarrow{-w} (q, \text{eps})\}$.

Zbiór terminali to oczywiście Σ . Natomiast nieterminal początkowy to S , dla którego mamy reguły: $S \rightarrow X_{q_0, f}^{s_0}$ dla każdego $f \in F$. Oprócz tego mamy następujące reguły.

Dla każdej tranzycji $\text{pop}(p, Y) \xrightarrow{-a} (q, \text{eps})$ robimy produkcję $X_{p,q}^Y \rightarrow a$. Natomiast dla każdej tranzycji $(p, Y) \xrightarrow{-a} (q, ZY)$ robimy produkcję $X_{p,r}^Y \rightarrow a X_{q,s}^Z X_{s,r}^Y$ dla dowolnych stanów r, s . Wystarczy zobaczyć, że ta gramatyka rzeczywiście generuje język L .

Zauważmy najpierw, że wszystko z języka może być rzeczywiście wygenerowane przez gramatykę. Patrzymy na bieg i aplikujemy odpowiednią tranzycję (na rysunku łatwiej zobaczyć). Z drugiej strony niech jakieś słowo będzie wygenerowane przez gramatykę. Zobaczmy na drzewo derywacji tego słowa w gramatyce. Patrząc na to drzewo od dołu budujemy kawałki biegów, które się dobrze składają. W korzeniu dostajemy bieg automatu ze stosem taki jak powinien być.
c.n.d.

Własności domknięcia

Zobaczmy jakie są własności domknięcia języków bezkontekstowych.

Po pierwsze są one zamknięte na **sumę**. Niech $L_1 = L(G_1)$ i $L_2 = L(G_2)$. Łatwo możemy zrobić G taką, że $L(G) = L_1 \cup L_2$. Po prostu robimy $S \rightarrow S_1$, $S \rightarrow S_2$, a resztę tak jak w gramatykach G_1 i G_2 .

Po drugie **nie** są one zamknięte na **przecięcie**. Łatwo pokazać, że $\{a^n b^n c^k \mid n, k \in \mathbb{N}\}$ oraz $\{a^k b^n c^n \mid n, k \in \mathbb{N}\}$ są bezkontekstowe. Natomiast ich przecięcie to $\{a^n b^n c^n \mid n \in \mathbb{N}\}$ co jak wiemy nie jest bezkontekstowe. Z tego wynika też, że języki bezkontekstowe **nie** są zamknięte na **dopełnienie**.

Warto natomiast zauważyć, że języki bezkontekstowe są zamknięte na **przecięcie z językami regularnymi**. Po prostu robimy produkt automatu ze stosem rozpoznającego dany język bezkontekstowy z automatem skończonym rozpoznającym dany język regularny. Jako wynik dostajemy znów automat ze stosem.

Deterministyczne automaty ze stosem

Dla automatów skończonych automaty deterministyczne i niedeterministyczne definiują tę samą klasę. **Pytanie** do sali: czy dla automatów ze stosem będzie tak samo? Nie, nie będzie tak samo. Trzeba najpierw powiedzieć dokładnie co to znaczy, że automat ze stosem jest deterministyczny. Bo teraz mamy też epsilony, musimy powiedzieć co znaczy determinizm przy ich obecności. Mówimy, że automat jest deterministyczny jeśli z danej konfiguracji albo da się wyjść epsilonem, ale niczym innym albo da się wyjść po każdej literze na tylko jeden sposób.

Okazuje się w szczególności, że deterministyczne automaty ze stosem są zamknięte na dopełnienie.

Tw.

Jeśli L jest rozpoznawalny przez deterministyczny automat bezkontekstowy, to dopełnienie L też.

Dowód (idea)

Gdyby nie było epsilonów to łatwo - po prostu zmieniamy stany akceptujące na nieakceptujące. Tu musimy trochę popracować z epsilonami.

Najpierw modyfikujemy nasz automat tak, żeby spełniał 3 warunki:

- symbol początkowy nie jest nigdy zdejmowany ani wkładany na stos
- zawsze jest przejście albo po epsilonie albo po każdej literze (i nic więcej)
- zawsze da się przeczytać całe słowo (nie ma pętli po epsilonie)

To się da dość łatwo zrobić (dodajemy nowy symbol, dodajemy brakujące przejścia, które nic nie robią, wykrywamy pętle po epsilonie). Teraz robimy automat dla dopełnienia. Robimy $Q' = Q \times \{0, 1, 2\}$. Druga współrzędna pamięta, czy dla danego słowa był już stan akceptujący, czy nie. Wartość 0 mówi, że stanu akceptującego nie było i nie będzie (bo zamierzamy właśnie przeczytać literę), wartość 1, że był stan akceptujący, a wartość 2, że nie było, ale czytamy jeszcze epsilony. Dodajemy opcję przejścia z 2 do 0 jeśli widzimy, że nie da się dalej iść epsilonami. Stan początkowy to $(q_0, 2)$, a końcowy to $(f, 0)$. To akceptuje dopełnienie.
c.n.d.

Własności domknięcia dla deterministycznych języków bezkontekstowych

Jak widzimy są zamknięte na dopełnienie.

Nie są też zamknięte na przecięcie, bo $\{a^i b^i c^j\}$ przecięte z $\{a^i b^j c^i\}$ daje $\{a^i b^i c^i\}$, które nie jest bezkontekstowe. Natomiast dwa poprzednie są.

W związku z tym nie są też zamknięte na sumę, bo inaczej byłyby na przecięcie też (bo na dopełnienie są).

Wykład 9

Problemy decyzyjne

Co chcemy wiedzieć o danej klasie języków? Chcemy znać:

- różne sposoby definiowania (tu gramatyki i automaty ze stosem)
- kiedy język jest w danej klasie (lemat o pompowaniu, konstruowania gramatyki lub automatu)
- jakie są własności domknięcia (znamy podstawowe)
- podstawowe problemy decyzyjne (teraz zrobimy)

Dla regularnych pytaliśmy podobnie, teraz zapytamy o:

- czy słowo w należy do języka L
- czy język L pusty
- czy dwa języki K i L są równe
- czy K zawiera się w L

Dla pierwszego problemu pokażemy zaraz algorytm, tzn. algorytm CYK (Cocke, Younger, Kasami), będzie działał w $O(n^3)$ (da się szybciej, w czasie $O(n^\omega)$, gdzie jest to najlepszy czas mnożenia dwóch macierzy).

Czy L jest pusty? Po prostu trzeba to sprawdzić na gramatyce. Patrzymy które zmienne mogą wyprodukować literę i patrzymy, czy początkowy symbol może przejść na taką zmienną. W czasie wielomianowym (być może liniowym).

Teraz okazuje się, że równość dwóch języków bezkontekstowych jest problemem nierozstrzygalnym, czyli nie ma dla niego algorytmu. Tym bardziej zawieranie. Wrócimy do tego przy okazji maszyn Turinga, bo wtedy będziemy umieli precyzyjnie dowodzić takie rzeczy.

Co ciekawe Senizergues w 1997 roku pokazał, dla dwóch języków akceptowanych przez deterministyczne automaty ze stosem równość jest rozstrzygalna (istnieje algorytm). To miało wiele stron (około 100 chyba). Teraz jest lepszy dowód Petra Jancara (może 2012 rok), około 15 stron, ale wciąż mocno skomplikowany. I nadal nie jest znana żadna rozsądna złożoność podczas gdy nie jest wykluczone (z tego co wiem), że problem jest rozwiązywalny w czasie wielomianowym.

Algorytm CYK

Mamy daną gramatykę G oraz słowo w , chcemy sprawdzić, czy w należy do języka $L(G)$.

Pomysł jest taki, żeby użyć programowania dynamicznego i sprawdzać dla coraz dłuższych infiksów słowa, czy ten infiks da się wyprodukować z danego terminala. Zrobimy więc tabelę trójwymiarową $T[i][j][X]$, która będzie zawierała 1 o ile słowo $w[i..j]$ należy do $L(X)$ (czyli da się wyprodukować z nieterminala X), a 0 w przeciwnym przypadku. Chcemy obliczyć $T[1][n][S]$, gdzie n to długość w , a S to symbol startowy. Zakładamy, że G jest w postaci normalnej Chomsky'ego (jeśli nie, to możemy ją do takiej przerobić). A więc najpierw liczymy $T[i][i][X]$. To jest prawda o ile tylko mamy w gramatyce produkcję $X \rightarrow a$, gdzie $a = w[i]$. Potem liczymy $T[i][j][X]$ dla coraz większych $j-i$. Po prostu $T[i][j][X]$ jeśli istnieje produkcja $X \rightarrow YZ$ oraz indeks k taki, że $i < k \leq j$ oraz $T[i][k-1][Y] = T[k][j][Z] = 1$. To się da obliczyć w czasie $O(j-i)$. A zatem cały algorytm działa w czasie $O(|G| n^3)$, bo każdą z $|G| n^2$ komórek da się wypełnić w czasie $O(n)$.

Uwaga: da się to zrobić w czasie $O(|G| n^\omega)$, (może razy $\text{polylog}(n)$) gdzie ω to minimalna liczba t , że mnożenie macierzy $n \times n$ da się zrobić w czasie n^ω . Widać, że liczenie $T[i][j][X]$ ma trochę wspólnego z mnożeniem macierzy.

Tw. Parikha

Twierdzenie Parikha to przydatny fakt o językach bezkontekstowych. Pokazuje intuicyjnie jakie języki bezkontekstowe możemy wygenerować jeśli nie interesuje nas kolejność liter, a jedynie to ile ich było. Bardzo ładny jest jego dowód i robię je tu również z tego powodu. Pokazuje, że żeby coś udowodnić o językach bezkontekstowych warto patrzeć na drzewa derywacji (podobnie jak to robiliśmy przy okazji lematu o pompowaniu) i jak to robić.

Najpierw kilka definicji. Niech $\Sigma = \{a_1, \dots, a_k\}$. Dla słowa w in Σ^* definiujemy jego obraz Parikha, który jest wektorem w N^k : $PI(w) = (\#a_1(w), \dots, \#a_k(w))$. PI do Parikh image. Ten wektor po prostu mówi ile jest których liter w słowie w . Obraz Parikha języka L to po prostu $PI(L) = \{PI(w) \mid w \in L\}$, czyli zbiór wektorów, które można uzyskać słowami z języka.

Twierdzenie Parikha będzie mówiło jakie obrazy Parikha mają języki bezkontekstowe.

Zbiór $S \subseteq N^k$ jest **liniowy** jeśli jest postaci $\{b + n_1 v_1 + \dots + n_d v_d \mid n_i \in N\}$, gdzie $b, v_1, \dots, v_d \in N^k$. Zbiór $S \subseteq N^k$ jest **semiliniowy** jeśli jest skończoną sumą zbiorów liniowych. Przykłady to:

- $\{(0,0) + n(1,1) \mid n \in N\}$ - przekątna
- $\{(0,0) + n(1,2) + m(2,1) \mid n, m \in N\}$ - skrzywiona krata
- $\{(10,10) + n(5,0) + m(0,5) + k(3,3) \mid n, m, k \in N\}$ - to jest taki dziwny zbiór
- suma wszystkich powyższych

Tw. Parikha

Obraz Parikha języka bezkontekstowego jest semiliniowy.

Przykłady:

1. $\{a^n b^n \mid n \in \mathbb{N}\}$ ma obraz Parikha $\{n(1, 1) \mid n \in \mathbb{N}\}$
2. $\{w \mid \#a(w) = \#b(w)\}$ ma taki sam obraz Parikha
3. $\{a^i b^j c^k \mid j = i+k\}$ ma obraz Parikha $\{k(1, 1, 0) + n(0, 1, 1) \mid k, n \in \mathbb{N}\}$
4. $L(a^*(bc)^*(a+c)^*)$ ma obraz Parikha $\{k(1, 0, 0) + n(0, 1, 1) + m(0, 0, 1) \mid k, n, m \in \mathbb{N}\}$

Uwaga: Język, który nie jest bezkontekstowy też może mieć semiliniowy obraz Parikha, np. $\{a^n b^n c^n \mid n \in \mathbb{N}\}$ oraz $(abc)^*$ mają ten sam obraz Parikha.

Dowód

Zakładamy, że mamy gramatykę w postaci Chomsky'ego.

Niech **pompa** to (prawie) drzewo derywacji z korzeniem o symbolu X, wszystkimi liśćmi z terminalami, ale jednym liściem z symbolem X. Jak mamy drzewo derywacji o symbolu X gdzieś w drzewie, to możemy je taką pompą **napompować**. Podobnie możemy pompować pompy.

Alfabet ALF(p) pompy p to zbiór wszystkich nieterminali w pompie p. Pompa jest **minimalna** jeśli nie da się jej uzyskać z mniejszej pompy o tym samym alfabetcie przez napompowanie. Podobnie drzewo jest **minimalne** jeśli nie da się go uzyskać z mniejszego drzewa o tym samym alfabetcie nieterminali. Zauważmy, że każde drzewo lub pompa minimalna ma głębokość nie większą niż n^2 (albo coś w tym stylu), bo większe drzewo można odpompować. A zatem jest skończenie wiele minimalnych drzew i minimalnych pomp. Teraz twierdzę, że każde drzewo derywacji można uzyskać z minimalnego drzewa derywacji o tym samym alfabetcie przez pompowanie minimalnymi pompami (o nie większym alfabetcie nieterminali). To w zasadzie widać. Co więcej jeśli pompujemy drzewa minimalne minimalnymi pompami, to osiągamy poprawne drzewo, każda z pomp o alfabetcie zawartym w S może być pompowana w minimalnym drzewie o alfabetcie S. A zatem drzewa = suma po minimalnych drzewach + pompowanie minimalnymi pompami o nie większych alfabetach. To daje natychmiast tw. Parikha.

c.n.d.

Uwaga: każdy semiliniowy zbiór da się uzyskać jako obraz Parikha języka bezkontekstowego, a nawet regularnego. Wystarczy np. dać $aab(abccc)^*(aaacc)^* + acc(ac)^*(abbbb)^*$, aby uzyskać $\{(1, 2, 0) + n(1, 1, 3) + k(3, 0, 2)\} \cup \{(1, 0, 2) + n(1, 0, 1) + k(1, 4, 0)\}$.

Wykład 10

Przykład zastosowania:

1. Czy język $\{a^{2^n} \mid n \in \mathbb{N}\}$ jest bezkontekstowy?
NIE, bo jego obraz Parikha nie jest semiliniowy. A dlaczego nie jest?
2. Czy język $\{a^n b^{n^2} \mid n \in \mathbb{N}\}$ jest bezkontekstowy?
NIE, z tego samego powodu.

Ciekawostka

Zbiory semiliniowe są ciekawymi zbiorami. Po pierwsze są zamknięte na operacje boolowskie:

sumę, przecięcie, dopełnienie. Zamknięcie na sumę jest oczywiste, ale na przecięcie i dopełnienie dalece nie jest. Po drugie zbiory semiliniowe to te same, które są opisywalne w tzw. arytmetyce Presburgera, czyli logice pierwszego rzędu nad liczbami naturalnymi, gdzie możemy używać dodawania. Przykładowa formuła: $\phi(x, y, z) := (x+z = y)$ opisuje zbiór semiliniowy z przykładu 3. Ten wynik też nie jest oczywisty, chociaż oba powyższe wyniki dadzą się zrobić na kilku stronach.

Maszyny Turinga

Zaczynamy teraz trzeci dział na tym przedmiocie: maszyny Turinga.

Pytanie: co to jest mechaniczny proces? Maszyna Turinga to po prostu jeden z równoważnych modeli obliczeniowych, którymi jesteśmy w stanie zamodelować każde możliwe obliczenie. Tak się złożyło, że teraz na całym świecie używa się maszyn Turinga właśnie, ale jest wiele równoważnych modeli. Zdefiniował ją tak Alan Turing (1912-1954) w roku 1936.

Jest wiele różnych, niewiele się różniących wersji maszyn Turinga, my teraz zajmiemy się jednym konkretnym. Maszyna Turinga to taki bardziej skomplikowany automat dwukierunkowy. Tyle, że automat dwukierunkowy mógł chodzić tylko po słowie (w lewo, w prawo i nie ruszać się) i dodatkowo nic nie mógł pisać, a tylko czytać i w zależności od tego robił ruchy. Natomiast maszyna Turinga będzie chodziła po taśmie, która jest nieskończona w obie strony. Na początku jest na niej napisane słowo wejściowe i maszyna zaczyna na lewym jego końcu. Dodatkowo maszyna może nie tylko czytać, ale też pisać.

Miejsca, które nie są zapisane literami słowa wejściowego w zapełnione są tzw. blankami, czyli specjalnymi literami oznaczającymi puste, nieodwiedzone miejsce. Maszyna ma też swój stan. W każdym ruchu maszyna patrzy na swój stan i literę napisanym na polu, na którym stoi. W zależności od tego zmienia stan, zapisuje na miejscu, na którym stoi jakąś literę (potencjalnie inną niż była) i idzie w którymś kierunku (lewo, prawo, stoi). Niektóre stany są akceptujące, a inne odrzucające. Jeśli maszyna wejdzie w taki stan, to akceptuje/odrzuca słowo wejściowe.

Formalnie, niedeterministyczna maszyna Turinga składa się z:

- alfabetu wejściowego Σ
- zbioru stanów Q
- stanu początkowego q_0
- zbioru stanów akceptujących F
- zbioru stanów odrzucających R
- alfabetu taśmowego Γ , $\Sigma \subseteq \Gamma$
- symbolu B (blank) należącego do $\Gamma \setminus \Sigma$
- relacji przejścia: $\delta \subseteq Q \times \Gamma \times Q \times \Gamma \times \{-1, 0, 1\}$

Konfiguracja maszyny jest postaci wq , gdzie w to słowo przed głowicą, q to stan, a w to słowo zaczynające się na komórce, na której jest głowica. Ignorujemy blanki (możemy założyć, że maszyna nigdy nie pisze blanków).

Konfiguracja początkowa maszyny to q_0w , czyli głowica maszyny stoi na pierwszej literze w , a stan to q_0 .

Tranzycja (q_1, a_1, q_2, a_2, k) oznacza: jeśli jesteś w stanie q_1 i widzisz pod głowicą a_1 , to zmień stan na q_2 , zapisz pod głowicą a_2 i przesunij głowicę o k komórek w prawo.

Bieg maszyny może się zakończyć (akceptując słowo bądź odrzucając), albo nie zakończyć. Zauważmy, że są różne sposoby jak może się nie zakończyć: może wpaść w cykl, ale może też np. cały czas w prawo, albo też inaczej. Ale to dla nas nie gra specjalnej roli. Może też być tak, że w pewnym staniu nie ma ruchu po pewnej literze, wtedy bieg jest odrzucany (ale staramy się robić maszyny tak, żeby tak się nie działo, wtedy po prostu przekierowujemy do stanu odrzucającego).

Słowo jest akceptowane jeśli istnieje chociaż jeden bieg akceptujący (w szczególności mogą być też inne biegi odrzucające). Język maszyny Turinga to zbiór słów akceptowanych. Maszyna jest deterministyczna jeśli relacja przejścia jest funkcją z $(Q \times \Gamma) \rightarrow (Q \times \Gamma \times \{-1, 0, 1\})$. Mówimy, że maszyna deterministyczna jest totalna jeśli zatrzymuje się (wchodzi do stanu z $F \cup R$) dla każdego wejścia.

Taśma jest nieskończona, jednak w każdej chwili używana jest jedynie jej skończony kawałek. Więc nieskończonej taśmy nie należy traktować jako nieskończonej pamięci, a raczej jako dowolnie dużą pamięć, co już jest rozsądnym założeniem w rzeczywistym świecie (przy modelowaniu). Maszyny Turinga będą służyły za formalny model wszelkich możliwych obliczeń. Za jej pomocą można symulować każdy możliwy algorytm, ale tego nie widać od razu.

Przykłady

Spójrzmy na następujący przykład maszyny. Będzie ona rozpoznawać język palindromów nad $\{a, b\}$. Idea jest tak, że sprawdzi pierwszą literę, zapamięta ją w stanie i zamaże ją #, pójdzie na koniec i sprawdzi, czy ostatnia litera jest taka sama. Jeśli tak, to zamaże ją \$ i pójdzie na początek, gdzie powtórzy to samo.

$\Sigma = \{a, b\}$

Stany to: $Q = \{q_0, q_a, q_b, t_a, t_b, p, ok, \text{źle}\}$

q_0 - początkowy, ok - akceptujący, źle - odrzucający

Alfabet taśmowy $\Gamma = \{a, b, B, \#, \$\}$, # będzie służyć do zamazywania początku, a \$ końca (dałoby się to obsłużyć jednym symbolem, ale jest jaśniej co się dzieje).

Tranzycje:

$(q_0, a, q_a, \#, +1), (q_0, b, q_b, \#, +1), (q_0, B/\$, ok, B/\$, 0)$

$(q_a, a/b, q_a, a/b, +1), (q_a, B/\$, t_a, B/\$, -1)$

$(q_b, a/b, q_b, a/b, +1), (q_b, B/\$, t_b, B/\$, -1)$

$(t_a, a, p, \$, -1), (t_a, b, \text{źle}, b, 0), (t_a, \#, ok, \#, 0)$

$(t_b, b, p, \$, -1), (t_b, a, \text{źle}, a, 0), (t_b, \#, ok, \#, 0)$

$(p, a/b, p, a/b, -1), (p, \#, q_0, \#, +1)$

Teraz inny przykład, pokazujący, że maszyny Turinga mogą robić też bardziej skomplikowane rzeczy. Będzie ona rozpoznawać język $\{ww \mid w \in \Sigma^*\}$.

Idea jest taka, że sprawdzi pierwszą literę, zapamięta ją w stanie i zamaże ją #, pójdzie w prawo i zgadnie w pewnym momencie, że tu jest początek drugiej połówki. Sprawdzi, czy litera się zgadza. Jeśli tak, to zamaże ją # i wróci po pierwszej połówce. Potem będzie robiła w kółko podobne rzeczy sprawdzając, czy litery w pierwszej połówce i z drugiej są takie same. Jak się zorientuje, że przeczytała wszystko i wszystko się zgadzało, to zaakceptuje, jak nie, to odrzuci.

Stany to: $Q = \{q_1, q_{1a}, q_{1b}, p, q_2, q_{2a}, q_{2b}, ra, rb, r, s, ok, \text{źle}\}$.

q_1 - początkowy, ok - akceptujący, źle - odrzucający

Tranzycje:

$(q_1, a, q_{1a}, \#, +1)$, $(q_1, b, q_{1b}, \#, +1)$, $(q_1, B, ok, B, 0)$
 $(q_{1a}, a/b, q_{1a}, a/b, +1)$, $(q_{1a}, a, p, \#, -1)$, $(q_{1a}, B, \text{źle}, B, 0)$
 $(q_{1b}, a/b, q_{1b}, a/b, +1)$, $(q_{1b}, b, p, \#, -1)$, $(q_{1b}, B, \text{źle}, B, 0)$
 $(p, a/b, p, a/b, -1)$, $(p, \#, q_2, \#, +1)$
 $(q_2, a, q_{2a}, \#, +1)$, $(q_2, b, q_{2b}, \#, +1)$, $(q_2, \#, \text{źle}, \#, 0)$
 $(q_{2a}, a/b, q_{2a}, a/b, +1)$, $(q_{2a}, \#, ra, \#, +1)$
 $(q_{2b}, a/b, q_{2b}, a/b, +1)$, $(q_{2b}, \#, rb, \#, +1)$
 $(ra, \#, ra, \#, +1)$, $(ra, a, r, \#, -1)$, $(ra, b/B, \text{źle}, b/B, 0)$
 $(rb, \#, rb, \#, +1)$, $(rb, b, r, \#, -1)$, $(rb, a/B, \text{źle}, a/B, 0)$
 $(r, \#, r, \#, -1)$, $(r, a/b, p, a/b, -1)$, $(r, B, s, B, +1)$
 $(s, \#, s, \#, +1)$, $(s, a/b, \text{źle}, a/b, 0)$, $(s, B, ok, B, 0)$

Wydaje się, że ta maszyna faktycznie akceptuje język $\{ww \mid w \in \Sigma^*\}$, ale może gdzieś być błąd, łatwo go zrobić w opisie maszyny. Tutaj zrobiliśmy tak precyzyjny opis to po, żeby go raz zobaczyć. Zazwyczaj będziemy po prostu mówić słowami co robi maszyna.

Maszyny wielotaśmowe

Maszyny Turinga, jak się okaże, opisują wszystko co trzeba, ale jednak są bardzo prostym (niskopoziomowym) modelem obliczeń. Jak widać wyżej nawet akceptacja prostych języków (jak $\{ww \mid w \in \Sigma^*\}$) wymaga dość skomplikowanych maszyn. W związku z tym rozważa się nieraz ich uogólnienia, przy pomocy których nieco łatwiej opisać różne algorytmy. Rozważymy tutaj maszyny wielotaśmowe, które będą miały tę samą wyrażalność co maszyny Turinga (czyli będą akceptowały te same języki), ale będzie odrobinę łatwiej się nimi posługiwać.

Idea jest taka, że maszyna wielotaśmowa (nazywamy je wsadowymi) zamiast jednej taśmy ma wiele taśm. Na każdej taśmie jest jedna głowica, natomiast stan jest wspólny dla wszystkich taśm, tj. jeden dla całej maszyny. Jedną z tych taśm jest **wejściowa**, czyli można z niej tylko czytać, ale nic nie można tam zmieniać. Reszta taśm to tzw. taśmy **robocze**, na których można robić co się chce. Na początku na taśmie wejściowej jest słowo wejściowe, a wszystkie pozostałe są puste (głowice są na początku słowa i gdziekolwiek - bo to nie ma znaczenia, każde miejsce pustej taśmy jest takie samo).

Formalnie rzecz biorąc konfiguracja maszyny wsadowej o $k+1$ taśmach (wejściowej i k roboczych) jest postaci: $Q \times (\Gamma^* \times \Gamma^*)^{k+1}$, czyli stan oraz na każdej taśmie to, co po lewej i po (nieściśle) prawej od głowicy.

Natomiast relacja przejścia δ jest zawarta w $Q \times (\Gamma^{k+1}) \times Q \times (\Gamma^k) \times \{-1, 0, 1\}^k$.

Okazuje się, że maszyny wsadowe rozpoznają te same języki co maszyny jednotaśmowe.

Tw.

Maszyny jednotaśmowe i wsadowe rozpoznają te same języki.

Dowód

Łatwo zasymulować maszynę jednotaśmową na wsadowej dwutaśmowej. Po prostu najpierw kopiujemy taśmę wejściową na taśmę roboczą, a potem działamy tak jak jednotaśmowa na taśmie roboczej.

W drugą stronę, czyli symulacja maszyny wsadowej $k+1$ taśmowej na maszynie jednotaśmowej jest trochę trudniejsza, ale też nie bardzo. Powiedzmy, że maszyna wsadowa ma alfabet Γ . Wtedy maszyna jednotaśmowa, która ją symuluje będzie miała alfabet $(\Gamma \times \{\text{true}, \text{false}\})^{k+1}$ cup Σ . Czyli jedna komórka taśmy odpowiada na $k+1$ komórek maszyny wsadowej, każda na innej taśmie. Wartość true/false odpowiada temu, czy na danej taśmie głowica jest właśnie na tej komórce. Musimy pamiętać pozycje głowic, bo teraz mamy tylko jedną. Sumujemy z Σ , żeby słowo wejściowe mogło być nad tym samym alfabetem. Teraz powiedzmy sobie co robi maszyna jednotaśmowa. Najpierw zamienia $w = a_1 a_2 \dots a_k$ na $(a_1\text{-true}, B\text{-true}, \dots, B\text{-true}) (a_2\text{-false}, B\text{-false}, \dots, B\text{-false}) \dots (a_k\text{-false}, B\text{-false}, \dots, B\text{-false})$ i wraca na początek. Teraz symuluje zachowanie maszyny wielotaśmowej. Trzeba pokazać, że umie zasymulować jedną tranzycję, skupmy się na tym. Jak to zrobimy to już będzie jasne, że ciąg tranzycji też da się zasymulować. Żeby zasymulować tranzycję maszyny wsadowej trzeba wiedzieć jakie symbole znajdują się pod głowicami. Więc maszyna jednotaśmowa skanuje wejście z lewej do prawej i zapamiętuje w stanie te symbole. Jak przeczyta wszystkie to już wie która tranzycja powinna być zaaplikowana (albo która może, w przypadku maszyny niedeterministycznej). Wtedy wraca do początku, idzie znów od lewej do prawej i jeśli trafi na którejś taśmie na głowicę, to zmienia głowicę na tej taśmie zgodnie z tranzycją.

c.n.d.

Uwaga: zauważmy, że jeśli maszyna wsadowa była deterministyczna, to jednotaśmowa też, to nam się później przyda.

Maszyny deterministyczne

Pokażemy teraz, że maszyny deterministyczne mają tę samą siłę wyrazu co maszyny niedeterministyczne. To co prawda nie będzie oznaczało, że nie będziemy dalej rozważać maszyn niedeterministycznych, bo może się okazać, że do opisanego tego samego języka

wystarczy mniejsza maszyna niedeterministyczna niż deterministyczna (tak jak w przypadku automatów).

Tw.

Każdy język rozpoznawany przez niedeterministyczną maszynę wsadową jest również rozpoznawany przez pewną deterministyczną maszynę wsadową.

Dowód

Na poziomie intuicyjnym symulacja polega na tym, że maszyna deterministyczna przeszukuje drzewo ścieżek maszyny niedeterministycznej. Konkretnie rzecz biorąc rozważmy jednotaśmową maszynę M niedeterministyczną o K tranzycjach. Skonstruujemy deterministyczną maszynę wsadową o 3 taśmach (2 roboczych) M' taką, że $L(M') = L(M)$. Maszyna M' na pierwszej taśmie roboczej generuje coraz dłuższe ciągi liczb: $n_1, \dots, n_m$ in $\{1, \dots, K\}$, które mają odpowiadać obliczeniom maszyny M . Następnie maszyna M' sprawdza na drugiej taśmie roboczej czy to obliczenie jest faktycznie akceptującym obliczeniem M . Jeśli tak, to akceptuje, a jeśli nie, to generuje następny ciąg liczb na pierwszej maszynie itd. Jeśli M ma pewien bieg akceptujący, to M' też i vice versa.

Wykład 11

Hierarchia Chomsky'ego

Noam Chomsky, znany lingwista, w 1956 opisał hierarchię typów gramatyk, znaną pod pojęciem hierarchii Chomsky'ego. Na każdym poziomie klasie gramatyk odpowiada klasa języków i klasa automatów.

Gramatyki typu 3, reguły: $A \rightarrow a, A \rightarrow Ba$
Języki regularne
Automaty skończone

Gramatyki typu 2, reguły: $A \rightarrow \text{gamma}$
Języki bezkontekstowe
Automaty ze stosem

Gramatyki typu 1, reguły: $\alpha A \beta \rightarrow \alpha \text{gamma} \beta$ (gramatyki kontekstowe), gdzie α oraz β to dowolne ciągi nieterminali i terminali
Języki kontekstowe
Maszyny Turinga z pamięcią ograniczoną liniowo

Gramatyki typu 0, reguły: $\alpha \rightarrow \beta$ (bez ograniczeń)
Języki **rekurencyjnie przeliczalne** (albo języki **częściowo rozstrzygalne**)
Maszyny Turinga (niedeterministyczne, konieczne kończące się)

Ta hierarchia moim zdaniem nie jest szczególnie ważna, bardziej z powodów historycznych warto ją znać. Warto też rozumieć jak różne klasy mają się do siebie.

Języki i problemy rozstrzygalne

Mówimy, że maszyny Turinga (jedno/wielotaśmowe, niedeterministyczne/deterministyczne) mają być modelem programów. Chcemy jednak żeby każdy program się kończył no i działał w zdeterminowany sposób. Dlatego sensownym modelem dla programów są tak naprawdę deterministyczne maszyny Turinga, które dla każdego wejścia kończą bieg. Inny słowy mówimy, że są totalne, albo posiadają własność stopu.

Języki, które są rozpoznawane przez takie maszyny to języki **rekurencyjne** lub też języki **rozstrzygalne**. Nie jest prawdą, że każdy język rekurencyjnie przeliczalny (częściowo rozstrzygalny) jest też rekurencyjny (przeliczalny). Idea jest taka, że maszyna może działać w nieskończoność i może nie być maszyny, która jest jej równoważna, ale zawsze się kończy. W tej chwili tego nie pokażemy, bo w ogóle nie wiemy, czy jest jakiś język nierozstrzygalny.

Języki a problemy

Język to podzbiór Σ^* . Problem decyzyjny to taki, dla którego dla każdego wejścia mamy odpowiedzieć TAK lub NIE. Przykładowe problemy decyzyjne to:

- dany graf G , wierzchołki s, t , czy $s \rightarrow t$?
- dana liczba n , czy jest pierwsza?
- dany automat A , czy $L(A)$ jest pusty?

Zauważmy, że problem decyzyjny można też widzieć jako język, mianowicie jako zbiór tych wejść, dla których odpowiedź jest TAK. Czyli powiedzmy problem osiągalności w grafie to zbiór słów postaci $\text{opis}(G)\#s\$t$ takich, że z wierzchołka s da się osiągnąć wierzchołek t w grafie G . Od tej pory będziemy utożsamiali problemy decyzyjne z językami, warto się do tego przyzwyczaić.

Inne równoważne modele

Jest wiele modeli równoważnych maszynom Turinga:

- deterministyczne maszyny Turinga
- niedeterministyczne/deterministyczne maszyny wsadowej
- gramatyki typu 0 (bez dowodu)
- automaty wielostopowe (pewnie zobaczymy to jeszcze)
- automaty z kolejką
- automaty z wieloma licznikami
- obwody logiczne
- itp. itd.

Rozstrzygalność a częściowa rozstrzygalność

Jak mówiliśmy nie jest tak, że każdy problem częściowo rozstrzygalny jest rozstrzygalny. To samo w sobie jest bardzo ciekawym faktem: nie każdy problem da się rozwiązać! Pokażemy to poniżej.

Natomiast prawdą jest następujący fakt (nie jest bardzo ważny).

Fakt

Każdy problem L taki, że zarówno L jak i $\Sigma^* \setminus L$ są częściowo rozstrzygalne jest rozstrzygalny.

Dowód

Niech M_1 rozpoznaje L , a M_2 rozpoznaje $\Sigma^* \setminus L$, obie deterministyczne (bo możemy tak założyć, jak pokazaliśmy wcześniej). Robimy M , która rozpoznaje L i się kończy. M symuluje naraz M_1 i M_2 na wejściu. Jak M_1 zaakceptuje, to akceptuje, a jak M_2 zaakceptuje, to odrzuca. Dla każdego wejścia w maszyna M się zatrzyma, bo albo w L (wtedy M_1 się zatrzyma) albo w $\text{not } L$ (wtedy M_2 się zatrzyma).

Kodowanie i maszyna uniwersalna

Łatwo sobie wyobrazić, że każdy obiekt, z którym mamy do czynienia można zakodować jako ciąg liter. Np. graf o n wierzchołkach można zakodować jako macierz sąsiedztwa. A automat można zakodować jako kodowania Q , Σ , F , q_0 , δ po kolei. Podobnie można również zakodować maszynę Turinga w słowie.

Bardzo przydatnym pojęciem jest **maszyna uniwersalna**. Bierze się ona właśnie z obserwacji, że dla każdej maszyny Turinga M istnieje pewne słowo, które ją koduje (Σ , Γ , Q , F , ...), nazwijmy je $\text{kod}(M)$. A zatem możemy rozważyć język $L = \{(\text{kod}(M), w) \mid M \text{ akceptuje } w\}$. Łatwo zauważyć, że język L można rozpoznać maszyną Turinga M' . Ta maszyna po prostu symuluje M na słowie w . Taka maszyna M' nazywa się maszyną uniwersalną. Czyli każdy problem rozstrzygalny da się sprowadzić do uniwersalnego problemu, który to jest językiem maszyny uniwersalnej.

Nierozstrzygalność problemu stopu

Dość zaskakującym odkryciem jest to, że istnieją problemy, dla których nie ma żadnego algorytmu, który by je rozwiązywał, teraz to pokażemy. Takich problemów jest wbrew pozorom całkiem sporo. Najprostszym chyba dowodem na to jest następujący argument: problemów jest tyle ile podzbiorów Σ^* , czyli continuum. Natomiast algorytmów (czy też maszyn Turinga) jest $\aleph_0$, czyli mniej. Muszą więc istnieć problemy decyzyjne, dla których nie ma algorytmów. Nie jest to jednak specjalnie ciekawe, bo chcemy raczej wiedzieć, że istnieją problemy, które da się w skończony sposób opisać i nie ma dla nich algorytmów (mówimy, że są nierozstrzygalne). I to jak się okaże jest dość powszechne. Najbardziej chyba kanonicznym (i znanym) problemem nierozstrzygalnym jest tzw. **problem stopu**:

Dane: deterministyczna maszyna Turinga M , słowo w

Pytanie: czy M zatrzymuje się na w

Gdyby problem stopu był rozstrzygalny, to rozstrzygalny byłby też **problem należenia**:

Dane: deterministyczna maszyna Turinga M , słowo w

Pytanie: czy $w \in L(M)$

Zrobilibyśmy bowiem maszynę M' , która działa tak jak M na początku. Jak M akceptuje, to M' też, a jak M odrzuca, to M' się pętli. A więc M akceptuje $w \Leftrightarrow M'$ zatrzymuje się na w .

Pokażemy więc, że problem należenia jest nierozstrzygalny.

Najpierw przypomnijmy znaną zagadkę o fryzjerze w miasteczku X, który strzyże wszystkich, którzy nie strzygą się sami. Takiego fryzjera nazwiemy superfryzjerem. Pytanie brzmi: czy superfryzjer strzyże sam siebie - szybko dochodzimy do wniosku, że żadna opcja nie jest możliwa. Czyli superfryzjer nie może istnieć.

Dowód

Niech miasteczko X zamieszkane będzie przez kody maszyn Turinga.

Powiemy, że $\text{kod}(M_1)$ strzyże $\text{kod}(M_2)$ jeśli M_1 akceptuje $\text{kod}(M_2)$.

Przypuśćmy, że problem należenia jest rozstrzygalny, skonstruujemy superfryzjera w miasteczku X. Jeśli problem należenia jest rozstrzygalny, to znaczy, że istnieje maszyna N taka, że

$L(N) = \{(\text{kod}(M), w) \mid M \text{ akceptuje } w\}$. Rozważmy teraz nową maszynę F taką, że jeśli wejście nie jest kodem żadnej maszyny Turinga, to odrzuca, a jeśli wejście to $\text{kod}(M)$, to akceptuje jeśli M nie akceptuje $\text{kod}(M)$. Mając N możemy łatwo skonstruować F, po sprawdzeniu, czy wejście jest kodem symuluje N na parze $(\text{kod}(M), \text{kod}(M))$, a potem neguje odpowiedź. Czyli F akceptuje wejścia tych maszyn Turinga (strzyże je), które nie akceptują samych siebie (czyli nie strzygą się). Czyli F jest superfryzjerem w miasteczku X. Sprzeczność, bo superfryzjer nie może istnieć.

c.n.d.

Nierozstrzygalność ciąg dalszy

Przyjrzyjmy się jeszcze raz temu co zrobiliśmy powyżej. Pokazaliśmy tak naprawdę, że problem należenia jest nierozstrzygalny. Pokazaliśmy poza tym, że problem należenia da się przeformułować jako problem stopu. A więc gdyby problem stopu był rozstrzygalny, to problem należenia też. Zatem problem stopu musi być również nierozstrzygalny.

Tę taktykę będziemy stosować cały czas do pokazywania, że pewne inne problemy są nierozstrzygalne. Będziemy pokazywać, że gdyby one były rozstrzygalne, to problem stopu (albo inny nierozstrzygalny) też byłby rozstrzygalny.

Formalnie rzecz biorąc ujmijemy to w pojęcie **redukcji**.

Niech K to język zawarty w Σ^* , a L język zawarty w Γ^* .

Redukcja z K do L to jest funkcja **obliczalna** $f: \Sigma^* \rightarrow \Gamma^*$ taka, że

$$w \in K \Leftrightarrow f(w) \in L$$

Zauważmy, że jeśli L jest rozstrzygalny i istnieje redukcja z K do L, to K też jest rozstrzygalny.

Istotnie, robimy następująco. Jak dostajemy $w \in \Sigma^*$, to aplikujemy do niego f. Potem patrzymy na $f(w)$ i pytamy, czy należy do L. To jest rozstrzygalne. Ale wiemy, że $f(w) \in L \Leftrightarrow w \in K$, czyli odpowiedzieliśmy na pytanie, czy $w \in K$. Czyli K też rozstrzygalny.

A więc redukcja z problemu nierozstrzygalnego do naszego (który rozważamy) pokazuje, że nasz jest nierozstrzygalny.

Pokażemy teraz, że różne problemy są nierozstrzygalne:

1. pustość języka maszyny: dana maszyna M , czy $L(M)$ jest pusty
2. pustość przecięcia języków dwóch gramatyk: dane gramatyki bezkontekstowe G_1, G_2 , czy $L(G_1) \cap L(G_2)$ jest pusty
3. uniwersalność gramatyki: dana gramatyka bezkontekstowa G , czy $L(G) = \Sigma^*$
4. problem odpowiedniości Posta (PCP - Post Correspondence Problem): dane pary słów $(u_1, v_1), \dots, (u_n, v_n)$, czy istnieje taki ciąg indeksów $i_1, \dots, i_k$, że $u_{i_1} u_{i_2} \dots u_{i_k} = v_{i_1} v_{i_2} \dots v_{i_k}$

No to dowodzimy po kolei nierozstrzygalności.

1. Robimy redukcję z problemu należenia do problemu pustości języka. Dana maszyna M i słowo w , konstruujemy maszynę M' taką, że $w \in L(M) \Leftrightarrow L(M') \neq \emptyset$. M' będzie zależała od M i od w . Maszyna M' po prostu ignoruje wejściowe słowo, potem wpisuje słowo w na taśmę i symuluje M . Jeśli $w \in L(M)$, to $L(M') = \Sigma^*$, wpp. $L(M')$ jest pusty. Czyli to jest redukcja do problemu niepustości, gdybyśmy umieli rozstrzygać pustość, to należenie też byśmy rozstrzygnęli. Podobnie można pokazać zresztą, że problem pełności maszyny Turinga jest nierozstrzygalny (można by potem zanegować wejście).

Wykład 12

Kontynuujemy pokazywanie nierozstrzygalności różnych problemów.

2. Robimy redukcję z problemu pustości języka maszyny M . Dla M konstruujemy dwie gramatyki G_1, G_2 takie, że $L(M) = \emptyset \Leftrightarrow L(G_1) \cap L(G_2) = \emptyset$. Zrobimy G_1 i G_2 takie, że $L(G_1) \cap L(G_2)$ będzie równy dokładnie poprawnym kodowaniom biegów maszyny M . Teraz powiemy co to jest kodowanie biegu. To jest $\#c_1\#\text{rev}(c_2)\#\dots\#c_k\#$ (co druga konfiguracja odwrócona), gdzie wszystkie c_i mają równą długość. Alfabetem to $\Gamma \cup \Gamma \times Q \cup \#$, litera z Γ oznacza komórkę taśmy, a litera z $\Gamma \times Q$ oznacza komórkę taśmy nad którą jest głowica (wpisujemy tam stan). Zobaczmy, że łatwo zrobić gramatykę, która rozpoznaje $w\#\text{rev}(w)$. Można też zrobić gramatykę, która rozpoznaje $c_i\#\text{rev}(c_{i+1})$, gdzie c_{i+1} powstaje z c_i przez jeden krok maszyny. Podobnie z $\text{rev}(c_i)\#c_{i+1}$. Niech więc K to język pierwszego typu $(c_i\#\text{rev}(c_{i+1}))$, a L drugiego typu $(\text{rev}(c_i)\#c_{i+1})$. Robimy G_1 i G_2 tak, żeby

$$L(G_1) = \# (\Sigma \setminus \#)^* \# (L \setminus \#)^* \cup (\#K)^* \#$$

$$L(G_2) = (\#K)^* \# (\Sigma \setminus \#)^* \# \cup \# (\Sigma \setminus \#)^* \# (L \setminus \#)^* \# (\Sigma \setminus \#)^* \#$$

Pierwsze części odpowiadają za kodowania nieparzystej długości, a drugie parzystej.

Nietrudno zobaczyć, że faktycznie $L(G_1) \cap L(G_2)$ rozpoznaje to, co trzeba.

3. Robimy redukcję z problemu należenia do problemu pustości języka. Skonstruujemy gramatykę G taką, że $L(M) = \emptyset \Leftrightarrow L(G) = \Sigma^*$. Konkretnie rzecz biorąc $L(G)$ będzie zawierać wszystkie słowa oprócz poprawnych kodowań akceptujących biegów M . Zauważmy najpierw, że język $\text{dop}(w\#w)$ jest bezkontekstowy. Możemy podobnie

pokazać, że język $\text{dop}(w\#next(w))$ jest bezkontekstowy. A więc robimy teraz G takie, że

$$L(G) = (\Sigma^* \#)^* \text{dop}(w\#next(w)) (\# \Sigma^*)^* \cup ((\Sigma \backslash \#)^* \backslash w_0 B^*)$$

Łatwo zobaczyć, że rzeczywiście $L(G)$ to wszystko oprócz poprawnych kodowań akceptujących biegów M .

4. Teraz PCP. Najpierw przykład, 3 pary: (b, bbb) , $(babbb, ba)$, (ba, a) .

Rozwiązanie to $(2, 1, 1, 3)$, bo

$babbb \ b \ b \ ba = ba \ bbb \ bbb \ a$

Będziemy robili redukcję z problemu należenia do PCP. Najpierw jednak pokażemy, że jeśli umiemy rozwiązywać PCP, to też coś więcej, tzn. PCP z ustaleniem która para jest pierwsza w rozwiązaniu. Rozważmy PCP $(u_1, v_1), \dots, (u_n, v_n)$. Robimy teraz nową instancję PCP (na przykładzie): (b^*, b^*b^*b) , (b^*, b^*b^*b) , $(b^*a^*b^*b^*, b^*a)$, (b^*a^*, a) , $(\$, \$)$. Łatwo zauważyć, że nowe PCP ma rozwiązanie wtw stare miało rozwiązanie zaczynające się od pary (b, bbb) . Czyli PCP z ograniczeniem jest nie trudniejsze niż ogólne PCP. Teraz zredukujemy problem należenia do PCP z ograniczeniem. Zrobimy takie PCP, by rozwiązania reprezentowały tylko poprawne kody akceptujących obliczeń maszyny M na słowie w , tzn. $\$c_1\#c_2\#...c_k\$$. Najpierw przeróbmy naszą maszynę tak, że na samym końcu, jak już jest w stanie akceptującym, to zamazuje wszystko specjalnym symbolem S i przechodzi na początek do specjalnego stanu f , nigdy też w trakcie biegu ten S nie jest użyty (będzie nam łatwiej zrobić to PCP).

Idea:

$(\$, \$c_1\#)$, (a, a) , (B, B) , $(\#, \#)$, $(abc, a'b'c')$ jeśli $a'b'c'$ powstaje w jednym kroku maszyny z abc , $(\#f, \text{eps})$, (S, eps) , $(S\$, \$)$.

Nietrudno zauważyć, że rzeczywiście rozwiązania naszego PCP odpowiadają biegom maszyny.

Złożoność obliczeniowa

Maszyny Turinga przydają się nie tylko do uściślenia faktu, że istnieje jakiś problem. Dobrze nadają się też do sformalizowania pojęcia “szybki program”, dla różnych wersji szybkości. Zajmiemy się teraz wstępem do złożoności obliczeniowej, która zadaje sobie pytanie jak bardzo złożone są problemy, to znaczy jak “szybki” jest najszybszy program, który może je rozwiązać. Oczywiście jak zwykle dbamy tu o zgrubne oszacowanie szybkości, asymptotyczne itd.

Rozważmy jakąś funkcję $f: \mathbb{N} \rightarrow \mathbb{N}$ taką, że $f(n) \geq \log(n)$. Wtedy oznaczamy:

$\text{DTIME}(f)$ = problemy, które da się rozwiązać w **czasie** $f(n)$ maszyną **deterministyczną**

$\text{NTIME}(f)$ = problemy, które da się rozwiązać w **czasie** $f(n)$ maszyną **niedeterministyczną**

$\text{DSpace}(f)$ = problemy, które da się rozwiązać w **pamięci** $f(n)$ maszyną **deterministyczną**

$\text{NSpace}(f)$ = problemy, które da się rozwiązać w **pamięci** $f(n)$ maszyną **niedeterministyczną**

Czas działania maszyny to liczba kroków, a **pamięć** to liczba komórek, które zostały odwiedzone. Jasne jest, że czas jest zawsze nie mniejszy (asymptotycznie) niż pamięć, bo jeśli

zrobiliśmy K kroków to odwiedziliśmy maksymalnie $K+|w|$ komórek pamięci (słowo wejściowe tu dodaje stały czynnik).

Definiujemy szereg bardzo ważnych klas przy użyciu powyższych metod:

$P (= PTIME) = \text{suma po } k \text{ naturalnych } DTIME(n^k)$

$NP = \text{suma po } k \text{ naturalnych } NTIME(n^k)$

$PSPACE = \text{suma po } k \text{ naturalnych } DSPACE(n^k)$

$NSPACE = \text{suma po } k \text{ naturalnych } NSPACE(n^k)$

$EXPTIME = \text{suma po } k \text{ naturalnych } DTIME(2^{\{n^k\}})$

$NEXPTIME = \text{suma po } k \text{ naturalnych } NTIME(2^{\{n^k\}})$

$L = DSPACE(\log(n))$

$NL = NSPACE(\log(n))$

Zachodzi $L \subseteq NL \subseteq P \subseteq NP \subseteq PSPACE = NPSPACE \subseteq EXPTIME \subseteq NEXPTIME \subseteq \dots$

Co ciekawe dzisiejsza nauka bardzo niewiele wie na temat ścisłego zawierania.

Wiadomo jedynie, że:

- L mniejsze od $PSPACE$
- P mniejsze od $EXPTIME$
- ogólnie, że dla f mniejszej istotnie od g zachodzi $DTIME(f)$ mniejsze od $DTIME(g)$ oraz $DSPACE(f)$ mniejsze od $DSPACE(g)$. To przy założeniu, że f i g rozsądne. Jeśli g większa od f^2 to już chyba starczy. Więcej na przedmiocie teoria złożoności.

P vs NP

Problem, czy P jest równe NP jest uznawany za obecnie najważniejszy problem informatyki teoretycznej. W 2000 roku Instytut Matematyczny Claya ogłosił 7 problemów uznanych przez nich za obecnie najważniejsze problemy matematyki. Za rozwiązanie każdego z nich fundowana jest nagroda miliona dolarów. A to oczywiście jedynie część chwały. Jest tam 6 problemów ściśle matematycznych i jeden z informatyki teoretycznej, właśnie, czy P jest równe NP . Zresztą moim zdaniem różnica między informatyką teoretyczną a matematyką jest umowna.

Przypomnijmy więc, NP to są te problemy, które są rozpoznawalne w czasie wielomianowym na niedeterministycznej maszynie Turinga. NP od Nondeterministic Polynomial, a na pewno nie od Non Polynomial, jak niektórzy mogą myśleć. Tutaj ważna jest intuicja tego niedeterminizmu, a chodzi o zgadywanie. Żeby to dokładnie zrozumieć przyjrzyjmy się najbardziej znanemu problemowi z klasy NP (o którym nie wiadomo, czy należy do P).

Jest to problem spełnialności. Dana jest formuła logiki zdaniowa używająca n zmiennych: $x_1, \dots, x_n$. Ta formuła może używać negacji, koniunkcji i alternatywy (wszystko pozostałe da się wyrazić).

Problem spełnialności (SAT od satisfiability) dla danej formuły zdaniowej ϕ pyta, czy istnieje takie wartościowanie zmiennych, że ϕ jest spełniona.

Jak pokazać, że problem SAT jest w klasie NP? Trzeba zaprojektować maszynę niedeterministyczną, która go rozwiązuje. Pomysł jest taki. Maszyna M na początku ma dwa przejścia, odpowiadające opcjom: $x_1 = 0$ oraz $x_1 = 1$. Potem ma dwa odpowiadające $x_2 = 0$ oraz $x_2 = 1$. Itd. aż na końcu ma dwa odpowiadające opcjom $x_n = 0$ lub $x_n = 1$. Powiedzmy, że maszyna zapisuje to wartościowanie na taśmie. Następnie maszyna sprawdza w czasie wielomianowym, czy to wartościowanie, które jest odpowiednie dla danego biegu jest poprawne. Jeśli jest to akceptuje, wpp. odrzuca. Zobaczmy, że jeśli istnieje jakieś wartościowanie spełniające, to istnieje też bieg mu odpowiadający i ten bieg będzie akceptujący. Wpp. nie ma dobrego biegu.

Jaki jest sens tej maszyny? Często się mówi, że zgaduje ona wartościowanie i potem sprawdza, czy dobrze zgadła. Tak się zazwyczaj myśli o problemach w NP: jest to zgadnięcie pewnego wielomianowego obiektu (początek biegu, który jest niedeterministyczny), a potem weryfikacja w czasie wielomianowym tego, czy to zgadnięcie było dobre.

Rozważmy inny znany problem w klasie NP. Jest to problem Komiwojażera. Dany jest graf skierowierowany (pełny) z wagami (liczbami naturalnymi) na krawędziach oraz pewna liczba T . Wagi określają czas przejazdu między wierzchołkami (miastami). Pytanie, czy komiwojażer (przedstawiciel firmy podróżujący w celu zdobywania klientów) może objechać wszystkie miasta w czasie co najwyżej T . Zakładamy przy tym, że nigdy nie odwiedza 2 razy tego samego miasta. Łatwo pokazać, że ten problem jest w NP. Po prostu maszyna zgaduje odpowiednią trasę komiwojażera, a potem weryfikuje, że jest dobra.

NP-zupełność

Okazuje się, że istnieją problemy "najtrudniejsze" w klasie NP. Takie problemy będziemy nazywać NP-zupełnymi. No bo każdy problem w P jest też w NP, warto zajmować się tymi problemami w NP, które są naprawdę trudne. Powiemy, że problem jest NP-zupełny jeśli jest w NP oraz jest NP-trudny. Natomiast problem L jest NP-trudny jeśli:

każdy problem w NP da się zredukować redukcją w czasie wielomianowym do L. To jest bardzo podobne do tego, co robiliśmy przy nierozstrzygalności. Mianowicie wyobraźmy sobie, że jest problem L, który jest NP-zupełny. Gdybyśmy pokazali, że L jest w P, to już łatwo pokazać, że $NP = P$. Weźmy bowiem dowolny problem K z klasy NP. Istnieje redukcja w czasie wielomianowym K do L. Przypomnijmy, że redukcje wielomianowa K do L to funkcja (teraz obliczalna w czasie wielomianowym) f taka, że

$$w \text{ in } K \Leftrightarrow f(w) \text{ in } L$$

A więc żeby stwierdzić, czy w należy do K liczymy $f(w)$ i patrzymy, czy należy do L. Ale na to pytanie da się znaleźć odpowiedź w czasie wielomianowym, bo założyliśmy, że L należy do P.

A zatem jeśli jakiś problem NP-zupełny jest w P, to $NP = P$. I to jest właśnie zaleta problemów NP-zupełnych. Apriori nie wiadomo wcale, czy jakiś problem NP-zupełny istnieje. Przecież tak jak dla każdej liczby naturalnej jest większa liczba naturalna to dla każdego problemu w NP mógłby istnieć od niego trudniejszy, ale wciąż w NP. Okazuje się, że jednak istnieją problemy

NP-zupełne. Co więcej jest ich całkiem sporo (jak już będziemy mieli jeden, to łatwo wskazać więcej sprowadzając ten jeden do nich, jak w nierozstrzygalności).

Jest to słynne twierdzenie Cooka z 1971 roku.

Tw. Cooka

Problem SAT jest NP-zupełny.

Dowód

Weźmy pewien problem L należący do NP.

Pokażemy jak dla słowa wejściowego w skonstruować taką instancję SATa $\phi(w, L)$, że $w \in L \Leftrightarrow \phi(w, L)$ jest spełnialna

Ponieważ L należy do NP, to istnieje pewna niedeterministyczna maszyna M_L , która rozpoznaje język L . A więc w należy do L wtw. gdy istnieje bieg akceptujący maszyny M_L na słowie w . Formuła $\phi(w, L)$ będzie właśnie kodować to, że istnieje taki bieg.

Niech słowo w ma długość n . Najpierw określmy jakie zmienne występowały w formule:

- $x(i, j, a)$: po i krokach maszyny na j -tej komórce taśmy jest litera a
- $x(i, q)$: po i krokach maszyna jest w stanie q
- $x(i, j)$: po i krokach głowica maszyny jest na j -tej komórce taśmy

Zauważmy, że skoro i jest ograniczone przez $O(n^k)$, to j też, a więc liczba zmiennych jest ograniczona przez $O(n^{2k})$. Można to ukonkretnić, np. $17n^9$, więc wiadomo ile tych zmiennych napisać. Teraz chcemy napisać formułę, która będzie spełnialna wtw. gdy istnieje poprawny bieg M_L na w :

- niesprzeczność (w każdej chwili dokładnie 1 stan, w każdej chwili głowica na dokładnie jednym miejscu, w każdej chwili w każdej komórce dokładnie jedna litera)
- konfiguracja początkowa (na odpowiednich komórkach odpowiednie litery, głowica na początku, stan początkowy)
- konfiguracja końcowa (istnieje i oraz q w F taki, że $x(i, q)$)
- poprawność przejść (jeśli nie ma gdzieś głowicy to się nic nie zmienia, jeśli jest, to istnieje tranzycja (czyli skończona alternatywa po tranzycjach) taka, że okolica się tak właśnie zmienia)

Widać, że dla słowa w skonstruowaliśmy w czasie wielomianowym formułę, która jest spełnialna wtw. gdy w należy do L . Czyli to jest faktycznie redukcja w czasie wielomianowym co kończy dowód.

c.n.d.

Pokażemy teraz, że kilka innych problemów jest NP-zupełnych. Tak naprawdę tych problemów jest nieskończenie wiele, a bardzo wiele naturalnych.

Najbardziej chyba znany to jest 3-SAT, czyli pewne wzmocnienie SATa.

Powiemy, że formuła zdaniowa jest w postaci 3-CNF (conjunctive normal form) jeśli jest koniunkcją **klauzul**:

$$\phi = C_1 \wedge C_2 \wedge \dots \wedge C_n,$$

gdzie każda klauzula jest alternatywą co najwyżej 3 **literałów**:

$$C_i = x_3 \vee \text{neg}(x_5) \vee x_6,$$

gdzie literały to zmienne lub ich negacje.

Przykład to:

$$(x_1 \vee \text{neg}(x_2) \vee x_3) \wedge (\text{neg}(x_3) \vee x_1 \vee x_4) \wedge (x_5 \vee \text{neg}(x_6) \vee x_2) \wedge (x_6 \vee \text{neg}(x_2) \vee \text{neg}(x_1))$$

Ta formuła jest spełnialna, gdyż wartościowanie: $x_1 = 1$, $x_2 = 1$, $x_6 = 1$, a pozostałe dowolne spełnia formułę.

Można tak zmodyfikować, dowód Tw. Cooka, żeby pokazać, że konstruujemy formułę w postaci CNF. Stąd mamy twierdzenie:

Tw

SAT dla formuł w postaci CNF jest NP-trudny.

Fakt

3-SAT jest NP-zupełny

Dowód

Łatwo pokazać, że 3-SAT jest w NP, po prostu zgadujemy wartościowanie spełniające i sprawdzamy, czy działa. Teraz żeby pokazać, że 3-SAT jest NP-trudny wystarczy zrobić redukcję z dowolnego problemu NP-trudnego. Wtedy dla dowolnego K w NP będziemy mieli redukcję: $K \rightarrow \text{inny NP-trudny} \rightarrow \text{nasz}$, czyli będziemy mieli redukcję do naszego. Zrobimy redukcję z SATa dla formuł w postaci CNF, więc z niego zrobimy redukcję.

Wystarczy pokazać, że formułę zdaniową ϕ w postaci CNF można przekształcić w czasie wielomianowym do formuły zdaniowej ψ takiej, że:

- ϕ spełnialna $\Leftrightarrow \psi$ spełnialna
- ψ jest w postaci 3-CNF

Robię następująco: jeśli mam klauzulę $(l_1 \vee l_2 \vee \dots \vee l_m)$ to robię nową zmienną x i przerabiam ją na koniunkcję dwóch klauzul:

$$(l_1 \vee l_2 \vee \dots \vee l_{(m-2)} \vee x) \wedge (\text{neg}(x) \vee l_{(m-1)} \vee l_m).$$

Łatwo zobaczyć, że ta koniunkcja jest spełnialna wtw. gdy wejściowa klauzula jest spełnialna.

Mogę tak robić produkując nowe klauzule wielkości 3 i zmniejszając moją do czasu gdy $m \leq 3$.

A więc sprowadzę w ten sposób do wielomianowo większej postaci 3-CNF.

c.n.d.

Wykład 13

Na początek musimy jeszcze udowodnić tw. Cooka, bo nie zrobiliśmy tego poprzednio.

Pokażmy jeszcze jedną przykładową redukcję. W problemie klikli mamy dany graf G i liczbę n (daną unarnie) i pytamy, czy w grafie G jest klika K_n jako podgraf indukowany.

Pokażemy następujący fakt.

Fakt

Problem klikli jest NP-zupełny.

Dowód

Łatwo pokazać bycie w NP: zgadujemy n wierzchołków i sprawdzamy, że to jest istotnie klika.

Teraz pokażemy NP-trudność, poprzez redukcję z 3-SATA.

Weźmy formułę: $\phi = C_1 \wedge C_2 \wedge \dots \wedge C_k$ o zmiennych $x_1, \dots, x_n$.

Stworzymy graf G i liczbę k taką, że ϕ spełnialna wtw. G ma klikę wielkości k .

Dla każdego literału w klauzuli tworzymy oddzielny wierzchołek. Mamy ich więc co najwyżej $3k$, czyli wielomianowo wiele. Łączymy dwa literały jeśli nie są w tej samej klauzuli oraz jeśli nie są tą samą zmienną, tylko raz zanegowaną, a raz nie. Twierdzimy, że w tym grafie jest klika wielkości $k \Leftrightarrow \phi$ jest spełnialna. Pokażemy to. Jeśli ϕ jest spełnialna, to istnieje wartościowanie takie, że w każdej klauzuli jest jakiś literał, który jest spełniony. Te literały z pewnością są połączone krawędziami, czyli tworzą klikę wielkości k . Z drugiej strony jeśli istnieje klika wielkości k , to możemy stworzyć wartościowanie, które przypisze tym zmiennym takie wartości, że są spełnione (nie ma tam x oraz $\neg(x)$), czyli ϕ jest spełnialna.

c.n.d.

Znane problem NP-zupełne

Można w podobny sposób pokazać, że jeszcze wiele problemów jest NP-zupełnych.

Znane przykłady to:

- czy w danym grafie istnieje cykl Hamiltona? (problem cyklu Hamiltona)
- problem komiwojażera
- dane n przedmiotów o wartościach w_i i objętościach o_i oraz plecak o całkowitej objętości O . Jak zapakować plecak, by miał maksymalną wartość? (problem plecakowy)
- dany graf, czy da się pokolorować jego wierzchołki na 3 kolory tak, by sąsiednie wierzchołki miały zawsze różne kolory? (3 kolorowanie)
- dany zbiór liczb całkowitych, czy istnieje jego niepusty podzbiór o sumie zero? (problem sumy podzbioru)
- itp. itd.

Równe, czy nie?

Jest bardzo wiele takich problemów i ludzie bardzo wiele myśleli nad tym, czy $P = NP$.

Większość naukowców uważa, że P jest różne od NP , ale są też tacy, którzy przychylają się bardziej do hipotezy, że $P = NP$, a my, ludzkość, po prostu bardzo mało wiemy o algorytmach.

Ja osobiście też uważam, że P jest różne od NP, a nie umiemy tego pokazać m.in. dlatego, że w ogóle prawie nic nie umiemy udowodnić jeśli chodzi o tę dziedzinę. Prawie że w ogólnie nie są znane techniki, które pokazują, że jakieś problemy nie są rozwiązywalne w czasie wielomianowym. W ogóle znanych jest bardzo mało dolnych ograniczeń w złożoności obliczeniowej (czyli dziedzinie, która zajmuje się klasyfikacją problemów względem ich skomplikowania, m.in. właśnie przyporządkowywania do różnych klas złożoności). Od lat 70-tych niewiele zrozumieliśmy. Jest pewien postęp, ale głównie w tym, że wiemy, że niektóre techniki nie pozwolą nam na pokazanie, że P jest różne od NP.

Próby obejścia

Są różne praktyczne podejścia do sprawy. Skoro jest wokół nas tyle problemów NP-trudnych to należy jakoś się do nich zabrać. Wybiły się następujące podejścia:

- algorytmy randomizowane (z dużym prawdopodobieństwem i szybko rozwiązujemy problem). Jednak jest niewiele problemów, które dają się zrobić w ten sposób, a nie znamy algorytmu w P.
- algorytmy aproksymacyjne (rozwiązujemy problem mniej więcej, np. znajdujemy nie najkrótszą ścieżkę komiwojażera, ale taką, która jest nie dłuższa niż najkrótsza + 1% trasy). W praktyce takie algorytmy wystarczą i okazuje się, że często dadzą się zrobić szybko.
- algorytmy wykładnicze o niskich podstawach. Np. zamiast 10^n robimy $(1,3)^n$. To pozwala na trochę większe n, ale też nie tak znowu duże. Według mojej (niewielkiej) wiedzy na ten temat to jest głównie teoretycznie rozważane.
- algorytmy heurystyczne. Działają zazwyczaj w miarę nieźle. Dość znane są tzw. SAT-solvery, czyli programy, które dla większości danych instancji SATa rozwiązują go szybko. To są po prostu heurystyki, które mówią - znajdź zmienną, która występuje w wielu miejscach i najpierw ją zgadnij, rozważaj klauzule z mniejszą liczbą literałów itd. Przy sprytnych metodach większość instancji SATa da się szybko zrobić.

Kryptologia

Jednak okazuje się, że nie wszystko da się szybko rozwiązać nawet heurystycznie.

Kryptologia opiera się na tym, że pewne rzeczy da się szybko zrobić w jedną stronę, a w drugą nie wiemy jak to zrobić szybko. Jednym z najbardziej popularnych takich problemów jest mnożenie dwóch liczb. Jak mamy dane dwie duże (powiedzmy 200-cyfrowe) liczby pierwsze, to możemy łatwo, w czasie $O(200^2)$ pomnożyć je. Jednak jak mamy ten iloczyn i chcemy go rozdzielić na czynniki pierwsze (dwa), to nie są znane żadne, nawet heurystyczne, metody, żeby to zrobić szybko. Nawet w sposób randomizowany, przybliżony (cokolwiek by to miało znaczyć w tym przypadku) itd. Czyli mówimy, że nie umiemy robić szybko faktoryzacji (tj. rozkładu na czynniki pierwsze). Dygresja: algorytmy na komputerach kwantowych zaczęły się od algorytmu Shora, który pokazał, że na (jeszcze nie zbudowanym w rozsądnej skali) komputerze kwantowym da się robić szybko faktoryzację (w czasie wielomianowym, przy użyciu losowości). Na tym się opiera kryptologia właśnie. Inny podobny problem to problem logarytmu dyskretnego. Jak mamy a, k, n, to umiemy obliczyć szybko $b = a^k \bmod n$. Ale jak mamy a, b, n,

to nie umiemy szybko obliczyć k takiego, że $b = a^k \bmod n$. Ten problem też jest często używany w kryptologii.

Zauważmy, że oba problemy są łatwo w NP. Dla faktoryzacji zgadujemy rozkład i łatwo sprawdzamy, że jest ok, a dla logarytmu dyskretnego zgadujemy k i też łatwo sprawdzamy. Nie wiadomo o tym żeby były NP-zupełne (zazwyczaj to się łatwo pokazuje, że duża szansa, że nie są).

Czyli gdyby $P = NP$, to byśmy mieli problem z kryptologią, bezpieczeństwem itd. Podobnie zresztą gdyby nauczono się budować komputery kwantowe, wtedy faktoryzacja i logarytm dyskretny byłyby rozwiązywalne łatwo. Powstała już dziedzina, która się nazywa kryptologia post-kwantowa i zajmuje się projektowaniem bezpiecznych szyfrów i protokołów na wypadek zbudowania komputera kwantowego.

PSPACE

Inną dość często rozważaną klasą jest PSPACE, czyli zbiór problemów, które dadzą się rozwiązać na maszynie Turinga w wielomianowej pamięci.

Ciekawe jest twierdzenie Savitcha, które mówi, że w tym wypadku niedeterminizm wcale nie dodaje siły wyrazu.

Tw. (Savitch)

$PSPACE = NPSPACE$

Pominiemy jego dowód.

W klasie PSPACE są również problemy zupełne, tj. PSPACE-zupełne. To są problemy, które są w PSPACE oraz są PSPACE-trudne, czyli każdy problem z PSPACE da się do nich sprowadzić za pomocą redukcji wielomianowej.

Typowy problem PSPACE-zupełny to QBF (Quantified Boolean Formulas), dane jest formuła postaci:

$$\text{exists_}x_1 \text{ forall_}x_2 \text{ exists_}x_3 \dots \text{forall_}x_{(n-1)} \text{ exists_}x_n \text{ phi}(x_1, \dots, x_n),$$
gdzie phi używa tylko zmiennych $x_1, \dots, x_n$. Pytanie, czy jest ona prawdziwa. Zobaczmy, że to jest trochę podobne do SATa, który pyta, czy formuła postaci:

$$\text{exists_}x_1 \text{ exists_}x_2 \text{ exists_}x_3 \dots \text{exists_}x_{(n-1)} \text{ exists_}x_n \text{ phi}(x_1, \dots, x_n)$$
jest prawdziwa.

Inny problem PSPACE-zupełny to uniwersalność automatu niedeterministycznego. Czyli: dany automat niedeterministyczny A , czy $L(A) = \Sigma^*$?

Wiele problemów PSPACE-zupełnych ma związek z grami. Na przykład ustalanie gracza posiadającego strategię wygrywającą w uogólnionych warcabach jest PSPACE-zupełne. Intuicja jest taka, że jeden gracz odpowiada za exists, a drugi za forall. Tzn. gracz pierwszy wygrywa

wtw. gdy istnieje jego ruch taki, że dla każdego ruchu przeciwnika istnieje jego ruch t., że dla każdego ruchu przeciwnika itd.

NL

Dość znana jest też klasa NL - problemy rozwiązywalne niedeterministycznymi maszynami w pamięci logarytmicznej. Tutaj rozważamy redukcje w pamięci logarytmicznej, a nie w czasie wielomianowym, bo te w czasie wielomianowym nie miałyby za bardzo sensu. Wtedy bowiem już sama redukcja mogłaby rozwiązać problem.

Najbardziej znanym problemem NL-zupełnym jest osiągalność w grafie: dany graf skierowany G i dwa jego wierzchołki s, t , pytanie, czy w grafie G z wierzchołka s da się osiągnąć wierzchołek t (tj. czy istnieje ścieżka).

Fajny jest algorytm osiągalności, który działa w NL. Po prostu zgadujemy ścieżkę, pamiętamy gdzie jesteśmy i idziemy krok dalej. Pamięć logarytmiczna nam wystarczy, bo pamiętamy tylko numer wierzchołka. Jednak musimy jakoś skończyć. Obserwacja jest taka, że jeśli da się osiągnąć t z s , to da się osiągnąć jakąś ścieżką o długości co najwyżej n . Czyli możemy kończyć przeszukiwanie po n krokach. Czyli robimy tak: zgadujemy kroki i liczymy je. Jak dojdziemy do n , to przerywamy i nie akceptujemy. Jeśli jest ścieżka z s do t , to ją znajdziemy, wpp. oczywiście nie. Używa pamięć jest logarytmiczna.

Używamy tu istotnie niedeterminizmu, bo zgadujemy ścieżkę. Innym dużym problemem otwartym jest, czy $L = NL$. Wydaje się, że raczej nie, ale naukowcy i tu nie są pewni.

Inne klasy złożoności

Wiadomo jedynie o niektórych klasach, że są różne. Tzn. nie wiadomo, czy P jest różne od $PSPACE$, ani czy L jest różne od P , ale wiadomo, że:

- L jest różne od $PSPACE$
- P jest różne od $EXPTIME$

W ogólności są twierdzenia o hierarchii (czasowej i pamięciowej), które mówią z grubsza, że dla rozsądnych funkcji f jeśli mamy funkcję $g > f^2$, to wtedy:

- $DTIME(f)$ ściśle zawarte w $DTIME(g)$
- $DSPACE(f)$ ściśle zawarte w $DSPACE(g)$.

To niewiele, ale zawsze coś.

A więc hierarchia jest nieskończona. Tzn. $PTIME$ istotnie zawarte w $EXPTIME$, a to w $2-EXPTIME$, a to w $3-EXPTIME$ itd. Są też problemy nierozwiązywalne w żadnym czasie $n-EXPTIME$, ale rozwiązywalne i nawet jeszcze trudniejsze. Ale to już temat na zupełnie inną opowieść. Trochę o tym na Złożoności Obliczeniowej.

Podsumowanie

Na tym przedmiocie uczyliśmy się o modelach obliczeń. Największy nacisk położyliśmy na trzy: automaty skończone, automaty ze stosem (i gramatyki) oraz maszyny Turinga. Jest też wiele innych, choć te trzy są rzeczywiście chyba najbardziej popularne.

Znając to, co wiemy możemy iść dalej, w kierunku precyzyjnego klasyfikowania problemów na trudniejsze i łatwiejsze. To zahacza trochę o filozofię i ogólnie zrozumienie świata. Dla zainteresowanych polecam np. te teksty:

- http://www.deltami.edu.pl/temat/matematyka/teoria_liczb/2016/12/27/Prawda_o_matematykach/
- http://www.deltami.edu.pl/temat/informatyka/algorytmy/2013/12/30/Na_granicy_mozliwosci/

Proszę o wypełnienie ankiet po tym przedmiocie. Szczególnie cenne będą dla mnie uwagi tekstowe. Dziękuję za cały semestr, to dla mnie było ciekawe doświadczenie, mam nadzieję, że dla Państwa również.